

G. H. Gittins

Proposing the SIRE Method: Extracting Informal Rules from Survey Data



Proposing the SIRE Method

Extracting Informal Rules from Survey Data: A Standardised Approach to Institutional Analysis Using Institutional Grammar

By

George Gittins
4737768

Master of Science

in Complex Systems Engineering and Management

at the Delft University of Technology

1st Supervisor:

Prof. dr. A. Ghorbani,

2nd Supervisor:

Prof. dr. T. Filatova,

Advisor:

Dr. S. Gil-Clavel

Key Terms: Data extraction, Decision trees, Institutional Grammar, Institutional social structures, Large language models, Method development, Parallel set diagram, Quantitative analysis, Research methodologies, Social structures, Socio-technical systems, Survey data.

Acknowledgements

First, I would like to thank my graduation committee: I am highly grateful to both my supervisors, Dr. Amineh Ghorbani, and Dr. Tatiana Filatova from TU Delft. In early 2023, I was fortunate to attend a guest lecture by Dr. Amineh Ghorbani, be intrigued by her work and have the opportunity to build upon it.

Amineh organised the committee with Dr. Tatiana Filatova and Dr. Sofia Gil Clavel who both helped to guide my work. Thank you to Dr. Tatiana Filatova for providing access to her resources, connections and consultation. A special thanks goes to Dr. Sofia Gil Clavel for the near-weekly support and advice made available to me during the many weeks of research.

I would like to thank Dr Sebastian Gabel, Dr Christopher Frantz and Thorid Wagenblast for giving me the time to speak with them about their research. My discussions with them helped to motivate, scope and develop my thesis.

This thesis would not be possible without the opportunity to work with the survey data collected thanks to the European Research Council project SCALAR (grant agreement no. 758014) funded under the European Union's Horizon 2020 Research and Innovation Program.

Finally, I would like to sincerely thank my family and friends. I would not be in this position if it weren't for you. There is nothing more important than a good second, third or fourth opinion.

Abstract

In the field of Institutional Analysis (IA), institutions describe the emergent rules that govern human behaviour and interaction. There are two key types: formal rules such as laws and policies, and informal rules such as norms and shared strategies. Evaluating the effectiveness of existing policies and developing new policies often requires a thorough understanding of informal rules. Traditional methods for identifying informal rules are predominantly qualitative, labour-intensive, and prone to bias, typically limited by small sample sizes. This thesis addresses the gap in standardised, quantitative approaches by developing a novel method to extract informal rules and social structures from survey data: Survey Informal Rule Extraction (SIRE).

The SIRE method involves preprocessing survey data, training decision trees to classify behaviours, extracting Institutional Grammar (IG) components, and generating structured Attribute, Deontic, Aim, Condition, Or else (ADICO) statements. The design process involves an iterative approach with three key steps: 1) exploring and comparing existing survey data from an institutional grammar lens, 2) evaluating different types of questions and identifying relationships in the data, and 3) testing various methods to analyse response data and convert questions into ADICO statements. These steps led to a systematic approach to understanding how informal rules shape human behaviour, supporting evidence-based policy design.

The effectiveness of SIRE is validated through three applications to datasets from the European Social Survey (ESS) and SCALAR. The first validation compares SIRE results with literature findings on household climate change adaptation, showing logical and comparable outcomes. The second validation demonstrates the method's utility in policy analysis, examining factors influencing the adoption of structural flood protection measures. The third validation showcased the method's potential as a strategic analytics tool for political parties, identifying conditions that increase the likelihood of voting for specific parties. The results demonstrate that the SIRE method has the potential to provide a clear representation of institutional rules, offering valuable insights for institutional and policy analysts.

While this thesis presents a novel method for extracting informal rules from survey data, limitations include potential generalizability issues due to the focused validation on specific datasets (SCALAR and ESS), AI-related inconsistencies in text processing, and methodological constraints such as the current inability to directly address open-ended questions or identify deontic components. Future work should focus on further validation across diverse contexts, integration with other analytical frameworks, and method automation. Improvements in data encoding, condition scope, and rule extraction could enhance the SIRE method's utility in policy development and behavioural analysis, addressing current limitations and expanding its applicability.

Contents

1. Introduction.....	7
1.1. Research Problem.....	7
1.2. Systems Perspective and Link to CoSEM Program.....	8
1.3. Research Approach.....	9
1.4. Outline.....	10
2. Research Methodology.....	11
2.1. Research Questions and Methods.....	11
2.3. Research Flow Diagram.....	13
2.2. Approach for Validating the Method.....	15
2.2.1. SCALAR.....	15
2.2.2. European Social Survey.....	16
2.2.3. Krefeld-Schwalb et al. (2024) Survey Data.....	16
3. Theoretical background.....	18
3.1. Institutional Grammar.....	18
3.2. Survey definitions.....	19
3.3. Decision trees.....	20
3.4. Outputs and visualisations.....	22
4. Literature Review.....	25
4.1. Challenges in Policy Analysis and Identifying Informal Rules.....	26
4.2. Utilisation of Surveys in Sociological Research.....	28
4.3. Applications of Institutional Grammar.....	29
4.4. Language Processing for Institutional Analysis.....	31
5. Proposing the SIRE Method: Extracting Informal Rules from Survey Data.....	33
5.1. Survey Data Examination.....	33
5.2. Selecting and Categorising Questions.....	34
5.3. Data Preparation and Preprocessing.....	35
5.4. Rewriting to ADICO Components.....	36
5.5. Creating Parallel Set Diagrams.....	37
5.6. Extracting Statements using Decision Trees.....	38
5.7. Visualise Extracted Statements.....	40
6. Validating the SIRE Method.....	43
6.1. Validation 1: Comparison with SCALAR Literature Results.....	43
6.2. Validation 2: Policy Analysis and Testing with SCALAR.....	51
6.3. Validation 3: Policy Analysis and Testing with ESS.....	58
6.4. Validation Summary.....	73
7. Reflections on the SIRE Method.....	74
7.1. Use Cases in Institutional Analysis Research.....	74
7.2. The Role of Institutional Grammar.....	75
7.3. Automating Steps with Generative Artificial Intelligence.....	77
8. Conclusions.....	79
8.1 Research Questions.....	79
8.2. Contributions.....	82
8.3. Limitations.....	84
8.4. Future work.....	85
References.....	87

Appendix.....	92
A. Literature Review Problem Identification Overview.....	92
A1. Literature on Climate Change Adaptation.....	92
A2. Literature Using the Institutional Network Analysis.....	94
B. SIRE Method Diagram.....	95
C. Standardised Python Scripts for the SIRE Method.....	96
C1. Survey Data Examination.....	96
C2. Selecting and Categorising Questions.....	96
C3. Data Preparation and Preprocessing.....	97
C4. Rewriting to ADICO components (LLM API Example).....	98
C5. Create Parallel Set Diagrams.....	100
C6. Extracting Statements Using Decision Trees.....	101
C7. Visualising Extracted Statements.....	105
D. Example Outputs of the SIRE Method Python Script.....	108
D1. Survey Data Examination and Categorising Questions.....	108
D2. Data Preparation and Preprocessing.....	110
D3. Rewriting to ADICO components (LLM API Example).....	110
D4. Create Parallel Set Diagrams.....	111
D5. Extracting Statements Using Decision Trees.....	111
D6. Visualising Extracted Statements.....	112
E. Limitations of Institutional Grammar and Mitigation Strategies.....	113

List of Abbreviations and Acronyms

ADICO	Attribute, Deontic, Aim, Condition, Or else
AI	Artificial Intelligence
BERT	Bi-directional Encoder Representations from Transformers
ESS	European Social Survey
GPT	Generative Pre-trained Transformer
IA	Institutional Analysis
IG	Institutional Grammar
INA	Institutional Network Analysis
LLM	Large Language Model
NLP	Natural Language Processing
SCALAR	Scaling up adaptation behaviour for models of climate change
SIRE	Survey Informal Rule Extraction

1. Introduction

1.1. Research Problem

In an increasingly complex and interconnected world, policymakers face unprecedented challenges in designing effective solutions to societal problems. The rapid pace of technological advancement, global economic shifts, and evolving social dynamics create a landscape where traditional approaches to policy-making often fall short. There is a growing need for evidence-based analyses that can inform strategic planning and institutional policy creation (Frantz, C. et al. 2021). However, a significant issue persists: the gap between formal policies and the informal rules that actually guide behaviour in practice. This disconnect often leads to ineffective policies and unintended consequences, showing the importance of developing more sophisticated methods for understanding and analysing both formal and informal institutions (Mesdaghi, B. et al. 2022). The following text explores this challenge and proposes a novel approach to bridging this gap, addressing a fundamental issue in the field of Institutional Analysis.

Institutional Analysis (IA) has emerged as a critical field of research to address these challenges, offering insights into the intricate web of rules, norms, and structures that shape human behaviour and societal outcomes. Within IA, institutions are emergent rules that shape human behaviour and interaction. These institutions can exist formally, as laws and policies, or informally, as norms and shared strategies. IA facilitates investigation into fundamental questions regarding the influences on decision-making, the structures of policymaking, and the organisation of public goods and services (IGRI, 2023). When designing new policies a common challenge is evaluating whether existing policies are effective and whether new policies will impact the behaviour of the population as intended. A typical approach is to compare the intended impact of a policy with the current informal rules in practice, human behaviours and routines. Figure 1.1. Provides a general overview of the IA and rule formalisation process.

Misalignment between informal behaviour and formal policy is used to inform policy changes. Example process:

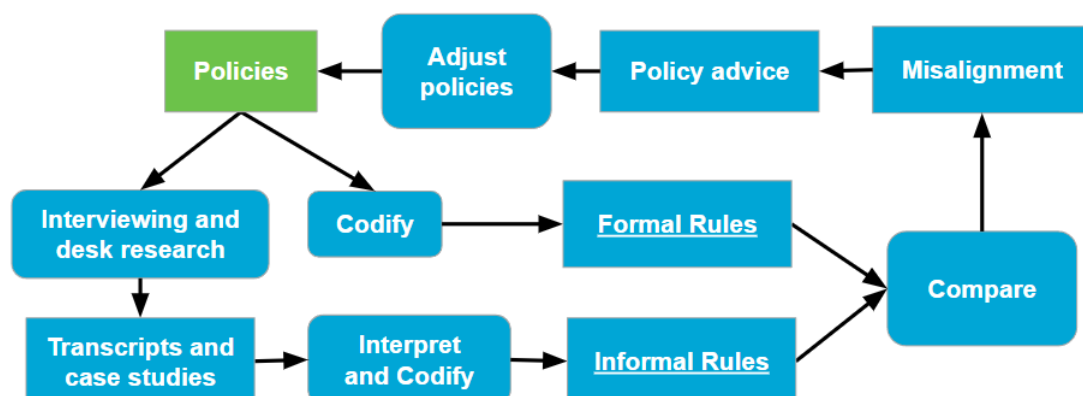


Figure 1.1: Diagram of the general IA process relevant to Identifying formal and informal rules.

Informal institutions are key components of IA, delineating behavioural patterns within demographics. Identifying formal rules involves reading and codifying legal documents and

conducting desk research, a structured and objective process relying on clearly defined methods. However, previous IA research has employed various methodologies to extract informal institutions such as formalising institutions from interview transcripts or conducting Institutional Network Analysis (INA) using desk research and interviews (Mesdaghi, B. 2020; Verheul, D. 2021; Juarez Pastor, L. 2022), a standardised approach focused on quantitative data does not exist. These studies are limited by labour-intensive methods that carry risks of bias and rely on small sample sizes, which do not objectively represent the population.

Institutional Grammar (IG) offers a structured approach to organising institutional content based on common features such as actions, actors, permissions, and constraints (Crawford, S. E. S. et al. 1995). By parsing institutional directives according to IG syntax, analysts can systematically analyse how institutions shape behaviour, providing a strong understanding of their scope and function. (Frantz, C. et al. 2021). Recent IA literature has introduced innovative uses of IG. For example, the IG Coder offers a specialised tool for encoding formal institutional statements, effectively addressing the complexities of institutional representations (Bognøy, J., 2021).

The goal for this thesis is to develop a standardised method for extracting informal rules and social structures from survey data. The IG framework serves as a robust foundation for this thesis. The research questions outlined in the next chapter will guide this study, aiming to enhance the objectivity of IA through a reproducible and scalable approach. The main research question of this thesis is:

What standard method can be developed to extract informal rules in use for institutional analysis from survey data?

In addition to answering the research questions, the following deliverables will be produced for this thesis:

- A step-by-step extraction method for informal rules from surveys.
- The code to reproduce this work on GitHub with advice for preparing survey data.

1.2. Systems Perspective and Link to CoSEM Program

This thesis embodies the Complex Systems Engineering and Management (CoSEM) program's holistic perspective on addressing complex socio-technical challenges. In the CoSEM context, navigating the convergence of innovations, regulatory frameworks, logistical hurdles, and behavioural dynamics is essential.

Key contributions include:

- **Standardised Methodology:** This thesis fills a critical gap in IA literature by developing a methodology to extract informal rules from survey data, for understanding practical institutional operations beyond formal regulations.
- **Integration of Diverse Methods:** Combining quantitative and qualitative techniques, such as decision trees and qualitative analysis, enhances clarity of social structures within institutional frameworks. This interdisciplinary approach is a core principle of CoSEM.

- **Systematic Approach to Institutions:** The designed method provides a representation of complex institutional interrelations, aligning with CoSEM's mission to prepare students for designing and managing intricate systems.
- **Practical Applications & Evidence-Based Policymaking:** Offering a robust method for extracting and using informal rules supports evidence-based policymaking and intervention design, translating theoretical frameworks into actionable strategies.
- **Interdisciplinary Collaboration:** The research showcases the CoSEM ethos by integrating technical and socio-political insights, promoting effective interdisciplinary collaboration.

This thesis advances IA methodologies and provides valuable insights for evidence-based policy and intervention design. It aligns closely with the CoSEM program's goals, equipping future systems engineers and managers to tackle the multifaceted challenges of modern society.

1.3. Research Approach

A Design Science Research Methodology has been used for this thesis. The Design Science process includes six steps: problem identification and motivation, definition of the objectives for a solution, design and development, demonstration, evaluation, and communication (Peppers, K. et al. 2014). The problem is first explored through desk research and discussions with Dr Amineh Ghorbani, Dr Sebastian Gabel, and Dr Christopher Frantz. This exploration helps to define the objectives and informs the iterative design process. The approach to producing the desired artefacts follows the systematic process outlined in Figure 1.2. Continuous evaluation and refinement take place throughout these stages to optimise the deliverables' efficacy and relevance.

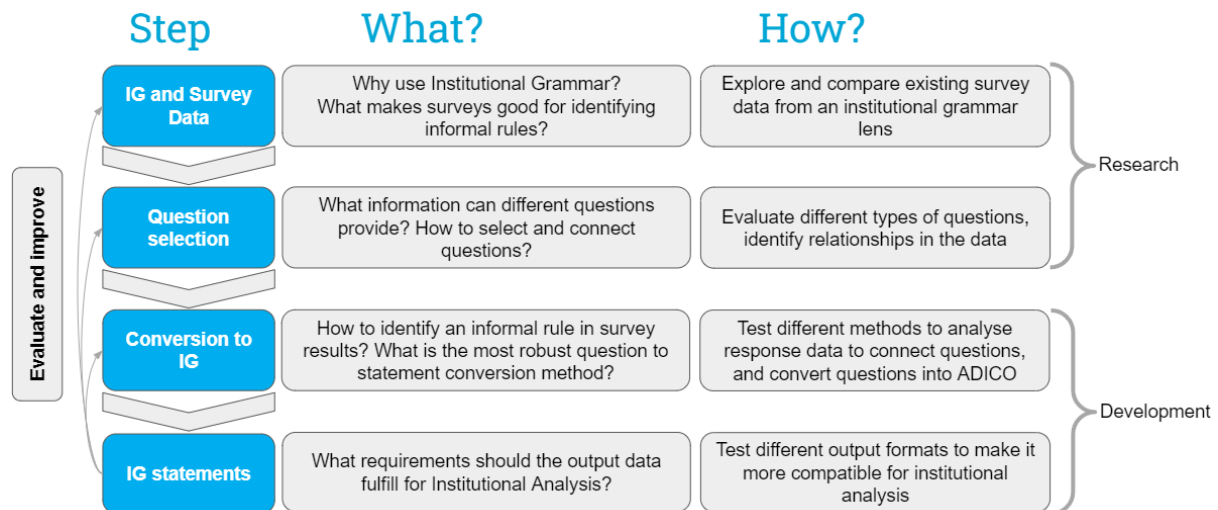


Figure 1.2: Research and development method for Master Thesis Method Design for Extracting Informal Rules from Survey Data.

The approach starts with survey data exploration, comparing existing surveys to identify characteristics conducive to extracting informal rules. Next, through question selection, various types of questions are evaluated to uncover relationships in the data, guiding the selection and connection of questions for analysis. Subsequently, different methods are

tested for converting questions into IG statements to identify the most robust conversion method. Finally, the study tests various approaches to making the output actionable and explores the optimal format for institutional analysts. The research methodology is explained in Chapter 2 and outlines the structured approach taken to develop and validate the Survey Informal Rule Extraction (SIRE) method.

1.4. Outline

The next chapter will discuss the research methodology, including the research questions and flow diagram, detailing the steps taken to develop a method to analyse survey data and extract meaningful IG components. Chapter 3 provides the theoretical background, covering Institutional Grammar, survey definitions, decision trees, and various outputs and visualisations relevant to the research. Chapter 4 conducts a literature review, exploring the challenges faced in policy analysis, the role of surveys in institutional analysis, applications of Institutional Grammar, and language processing techniques. Chapter 5 describes the SIRE method, explaining survey data examination, question categorisation, data preparation, rewriting questions as ADICO components, creating parallel set diagrams, and extracting statements using decision trees. Chapter 6 validates the SIRE method through policy analysis and testing with SCALAR and ESS data. Chapter 7 reflects on the SIRE method, discussing its use cases in institutional analysis, the role of Institutional Grammar, and the potential to automate steps with generative AI. Chapter 8 concludes the thesis by addressing the research questions, summarising contributions, discussing limitations, and suggesting future work.

2. Research Methodology

2.1. Research Questions and Methods

In this chapter the research methodology of this thesis is outlined. First, the overarching approach is described to explain how the main research question will be answered and how the desired artefacts will be produced. The main research question is:

What standard method can be developed to extract informal rules-in-use for institutional analysis from survey data?

The thesis employs the Design Science process (Peppers, K. et al. 2014), this process encompasses six stages: problem identification and motivation, defining solution objectives, design and development, demonstration, evaluation, and communication. The main research question and identified knowledge gaps guided the formulation of sub-questions, each addressing a specific aspect of the research problem. Each sub-question progresses towards building a robust method for extracting informal rules from survey data. The following sub-questions (SQs) have been derived from the main research question and the gaps and limitations discovered through the literature review. The sub-questions provide a structure for the steps taken to answer the main research question.

Sub-Question 1:

What challenges can institutional and policy analysts face when identifying institutional informal rules?

Sub-question 1 steers research towards identifying the obstacles institutional and policy analysts encounter when extracting informal rules. By understanding current methods and their limitations, and exploring the role of survey data, the goal is to pinpoint areas for improvement. This understanding will inform the design requirements for a more effective and standardised extraction method.

Sub-Question 2:

To what extent does Institutional Grammar align with survey data to identify and organise informal rules?

This sub-question evaluates the efficacy of Institutional Grammar (IG) in structuring and interpreting survey data. The research approach combines theoretical review with practical application to assess IG's potential for capturing and exploring the complexities of social structures and policies within survey responses.

An in-depth exploration of IG's structured approach reveals its suitability for organising institutional content. Key components of IG, such as Attributes, Aims, and Conditions, are found to map effectively onto survey questions, providing a robust framework for analysis. The study leverages existing datasets (SCALAR, ESS) to demonstrate IG's applicability in extracting meaningful informal rules from survey data. This practical assessment involves developing a method that uses IG to structure survey responses and identify social behaviours, thereby validating its utility in real-world contexts.

The findings support the use of IG as a valuable tool for organising and presenting human behavioural patterns derived from survey data. By examining current applications of IG by analysts and understanding the nature of survey data, this research contributes to improving IG's integration with survey-based studies, potentially enhancing its application in institutional analysis. This refined methodology provides a clear rationale for choosing IG, demonstrates its alignment with survey data, and highlights its potential for accurately capturing complex social structures and policies.

Sub-Question 3:

What are the possible approaches to connect survey results to Institutional Grammar concepts?

This sub-question explores methodologies that effectively align Institutional Grammar (IG) with survey data. The research aims to refine the extraction process of informal rules by identifying key IG components and developing systematic methods to map survey questions to these components. This process involves bridging theoretical constructs with empirical evidence and investigating data analysis techniques to uncover institutional social structures from survey responses. This assessment provides insights into the strengths and limitations of survey-based data collection for IG analysis. The research identifies and implements the following effective analysis methodologies:

1. **Literature Review and Practical Testing:** A comprehensive literature review establishes the theoretical foundation for linking IG with survey data. Concurrent practical tests on existing surveys identify successful practices and potential challenges, ensuring a balance between theory and application.
2. **Survey Design and IG Component Mapping:** The study tests various methods for identifying IG components within survey questions. This process leads to the selection of key tools and the refinement of the approach, optimising the alignment between survey design and IG structure.
3. **Data Structure Analysis:** An in-depth analysis of dataset properties, including variables, relationships, and patterns, provides crucial insights into effectively aligning survey data with IG concepts. This step ensures that the data structure supports the extraction of meaningful informal rules.

In the development of the informal rule extraction method, the potential for AI automation of each step in the process will be considered. While the priority is to explain and validate how each of these steps can be performed manually, exploring the potential for automation will provide insights into current and future advancements of the method.

Through these methodological steps, the research ensures that survey data is structured to facilitate the extraction of informal rules. The continuous testing of various analysis tools and refinement of the analytical approach contribute to the development of a robust methodology for connecting survey results to IG concepts. This process aims to produce the most effective combination of techniques, yielding insightful and actionable IG statements from survey data.

Sub-Question 4:

What requirements should the output data fulfil to ensure that the informal rules from surveys are usable for Institutional Analysis, consistent and well-communicated?

This sub-question addresses the quality assurance of extracted informal rules, with a focus on their usability and consistency for institutional analysts. The research aims to determine the practical applications of these outputs in Institutional Analysis (IA) and identify the most effective formats for presenting informal rules. Additionally, it explores visualisation techniques to optimise the communication of results, leveraging methods such as decision trees to enhance the informativeness and applicability of the generated statements. The study employs a comprehensive approach to identify the requirements for output data that is usable, consistent, and well-communicated:

1. **Review of Existing Techniques:** The research begins with a thorough examination of current data analysis techniques and feature extraction methods. This review informs the development of a tailored extraction method suitable for the specific dataset used in the study. By building on established practices, the method ensures a solid foundation for extracting meaningful informal rules.
2. **Evaluation of Potential Use Cases:** The effectiveness of the developed method is assessed by comparing the extracted structures to known IA literature. This comparative analysis serves a dual purpose: it validates the relevance of the extracted rules within the IA context and identifies areas for refinement. The process of evaluation and refinement ensures that the outputs are compatible with potential use cases of the method, enhancing its practical utility for institutional analysts.
3. **Visualisation and Presentation Techniques:** The study explores various visualisation techniques to enhance the clarity and interpretability of the extracted informal rules. This includes experimenting with different formats and graphical representations to optimise the communication of complex institutional structures.

The research implements an iterative approach to method development and refinement. This ongoing process allows for continuous adaptation of the extraction and presentation techniques based on feedback and emerging insights. The iterative nature of the methodology ensures that the method remains flexible and effective in meeting the evolving needs of IA practitioners. By addressing these key areas, the research aims to establish a set of comprehensive requirements for output data. These requirements ensure that the informal rules extracted from surveys are not only usable and consistent for IA but also effectively communicated. The resulting methodology provides a robust framework for generating clear, actionable insights from survey data, thereby enhancing the practical application of informal rules in institutional analysis.

2.3. Research Flow Diagram

The flow diagram presented in Figure 2.2 outlines the systematic approach employed in this master thesis. This diagram serves as a visual roadmap of the methods used to address each sub-question in the research, guiding the process of standardising and automating the extraction of institutional informal rules from survey data. Each stage of the research is detailed, starting with a background review of relevant literature on IG and survey analysis, followed by the exploration and operationalization of data extraction methods.

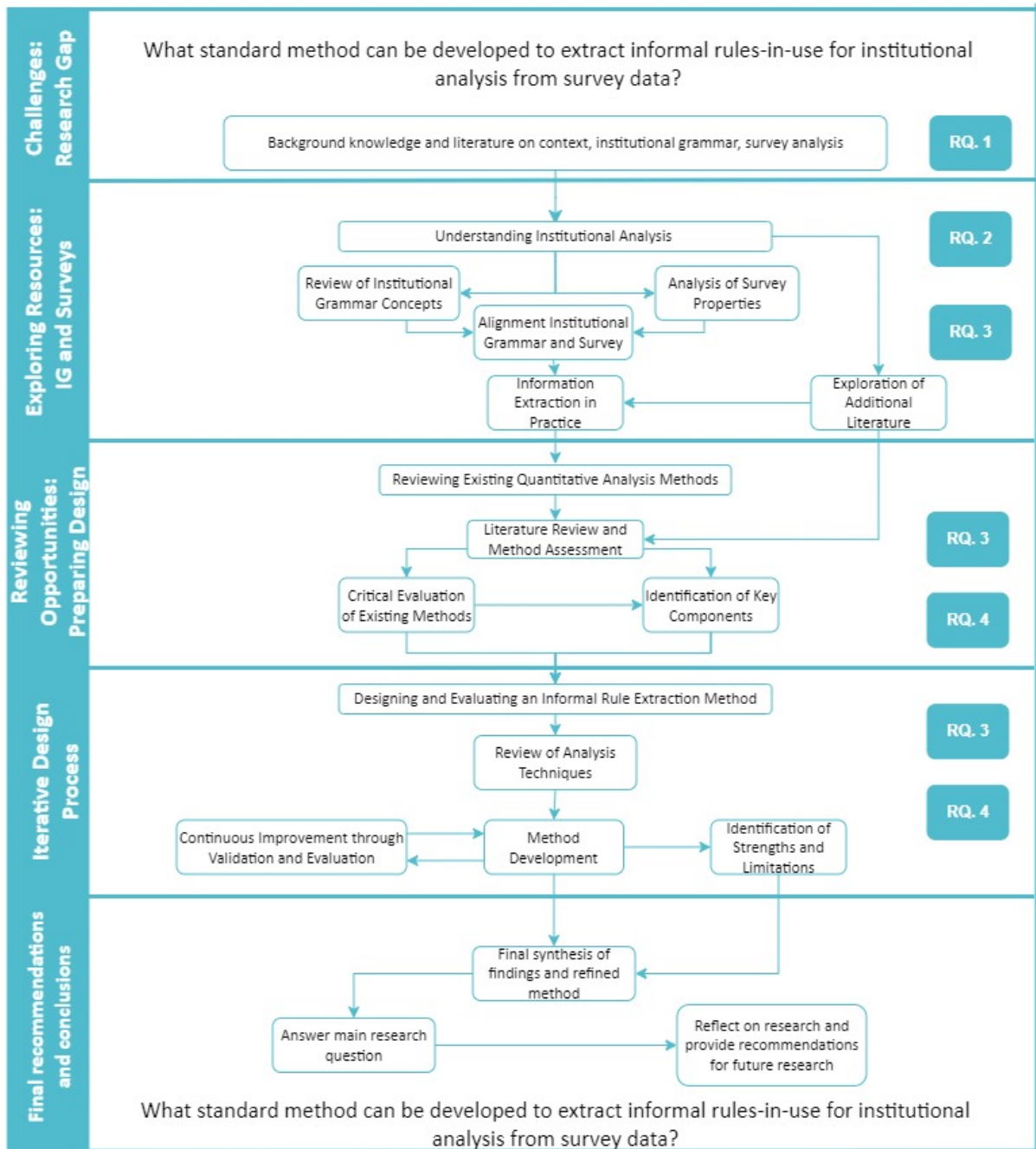


Figure 2.2: Research Flow Diagram for Master Thesis: Identifying Institutions and Social Behaviour: A Method for Extracting Institutional Informal Rules from Surveys

Key components in Figure 2.2 include the alignment of IG with survey data, the critical evaluation of existing quantitative analysis methods, and the iterative design process of an informal rule extraction method. The process culminates in a continuous loop of evaluation and refinement, ensuring that the method developed is both effective and applicable to a variety of survey contexts. This structured approach not only addresses the research questions methodically but also enhances the reliability and validity of the findings, providing a framework for future research in IA.

2.2. Approach for Validating the Method

In developing the informal rule extraction method, three main surveys and their use cases have been examined: the SCALAR survey, the European Social Survey (ESS), and the Krefeld-Schwalb et al. (2024) survey data. This section describes and contextualises each of these surveys. Additionally, The SCALAR survey and ESS data will be used to test and validate the extraction method, as well as show potential applications of the method. Three types of validations for the method can be applied:

1. **General Outcomes:** When applying the method to each dataset, the outputs are assessed for logicity, cleanliness, and explainability. This involves examining whether the method and outputs behave as expected and produce understandable results.
2. **Comparison with Case Studies:** The results obtained from the method will be compared with those from existing studies identified in the literature review that have used the same datasets. Validating the method involves demonstrating that it can produce similar conclusions to those found in the literature, ensuring the reliability and robustness of the approach.
3. **Policy and Informal Rule Comparison:** To showcase a practical application of the method, existing or proposed policies linked to the survey questions will be identified. By comparing the extracted informal rules with these policies, the effectiveness of the policies in influencing the intended behaviours can be evaluated. This comparison will help illustrate how the method can be used to assess and refine policy interventions.

Through these tests, the informal rule extraction method will be validated, demonstrating its applicability and accuracy in analysing survey data and contributing to IA. The detailed examination of these surveys and their use cases will provide a better understanding of the method's potential and limitations.

2.2.1. SCALAR

SCALAR is a panel longitudinal survey that contains responses from households located in large, coastal urban centres in the United States (Houston, New Orleans and Miami greater areas), the Netherlands (Rotterdam greater area and the province of Zeeland), China (Shanghai greater area), and Indonesia (Jakarta greater area, other cities in Java). The panel surveys conducted in five waves between 2020-2023 are focused on soliciting information on households' various factors influencing individual CCA behaviour surrounding private adaptation to climate change (Noll, B. et al. 2022). For the validation process the responses from the first wave in the Netherlands have been used, this data contains responses from 1251 participants.

The questionnaire was not specifically designed for IG statement extraction. The dataset has been used for various analyses (Wagenblast, T. 2022; Lechner J. 2022; Noll, B. et al. 2022). The dataset contains questions and responses that can be represented as informal rules. The “Or Else” element is not present in the questions and answers of the questionnaire. Therefore, formal rules cannot be directly extracted from the dataset as these require desk research on the currently active policies in the focus areas (Lechner J. 2022). However, the

results of the dataset analysis will be interesting because they could be used to establish which existing informal strategies should be encouraged or discouraged by changing the formal rules in place (Breza, E. et al. 2014).

2.2.2. European Social Survey

The ESS is a pan-European research infrastructure designed to collect and provide freely accessible data for academics, policymakers, civil society, and the wider public. Established to maintain high-quality standards in cross-national surveys, the ESS was granted European Research Infrastructure Consortium status in 2013. This achievement reflects the ESS's commitment to advancing methodology and setting benchmarks for cross-national surveys (European Social Survey, 2024). The ESS collects cross-sectional data, which means it gathers information at specific points in time from a wide range of European countries (Longford, N. T. 2008). This data provides a snapshot of various social, cultural, and political phenomena, allowing researchers to analyse diverse European issues. Unlike longitudinal data, which tracks the same individuals over an extended period, cross-sectional data focuses on the specific time frame when the data is collected, without necessarily following the same subjects in future waves or surveys. For the validation process the data from ESS10 edition 3.2 has been used. The survey was performed in 2020 and includes responses from 37611 participants, of which 1470 are from the Netherlands.

The ESS methodology emphasises rigour and consistency; this focus on quality has made it a valuable resource for researchers seeking to understand European social trends. Two notable studies utilising ESS data are by Pavlova et al. (2023) and Chueri et al. (2023). Pavlova et al. (2023) used the ESS to explore the relationship between voluntary participation in nonpolitical volunteering, involvement in political matters and well-being across different age groups and European countries. The study's extensive list of control variables and pre-registered statistical analyses added to its validity, though the cross-sectional design limited causal inference. Chueri et al. (2023) examined how descriptive and substantive representation affects women's voting patterns for populist radical right parties (PRRPs), employing ESS data and conditional logit models.

2.2.3. Krefeld-Schwalb et al. (2024) Survey Data

The survey data analysed by Krefeld-Schwalb et al. (2024) was collected independently across large samples in three countries: China, the United States, and The Netherlands. This geographical scope offers a broad view of sustainable behaviours and associated motives across diverse populations. The survey, conducted in English in the United States and The Netherlands, and translated into Mandarin for China, included questions about the respondents' living situations, demographic variables, and access to cars and public transport. This survey will not be used for validation but their methodology and data formatting has been examined to inform the design of the SIRE method.

The survey demonstrates the benefits of designing with a specific research focus, which can inspire this thesis' attempt to standardise survey design for institutional analysts. Their tailored approach with customised questions and randomised order, enabled them to gather detailed information in alignment with their research objectives. By creating a survey that targets particular behaviours and collects supporting data on motives, risks, and

preferences, they achieved a nuanced understanding of sustainable practices across multiple contexts. This customisation and attention to detail are valuable lessons for developing standardised survey designs that can accurately capture the elements of IG and provide a robust framework for IA.

This survey data is cross-sectional, providing a snapshot of the respondents' behaviours and motivations at a single point in time. This cross-sectional approach is suitable for exploring current trends and relationships between various factors and sustainable behaviours, although it does not allow for analysing changes or developments over time (Longford, N. T. 2008). The extensive data collected across different countries and the large range of questions offer valuable insights into sustainable behaviour patterns and underlying motives.

3. Theoretical background

This chapter outlines the theoretical framework and methodological tools that underpin this research. It begins with an in-depth exploration of Institutional Grammar (IG), elucidating its key components and establishing its critical role in this study. The discussion then transitions to an examination of survey methodologies, focusing on the nature of information gleaned from survey questions and their potential alignment with IG components. Subsequently, the chapter explores the application of decision trees as a means to uncover relationships within survey results. The final section reviews potential data outputs and visualisation techniques that can effectively communicate findings. This comprehensive theoretical foundation serves as a guide for developing a standardised method to extract and analyse informal rules from survey data, ensuring the robustness and relevance of results for Institutional Analysis (IA).

3.1. Institutional Grammar

IG is a tool for analysing institutions by decomposing them into base components, providing a framework for understanding social systems and structures (Bognø, J., 2021). IG uses the Attribute, Deontic, Aim, Condition, Or else (ADICO) syntax to standardise institutional statements:

- **Attribute (A):** The entities acting (e.g., "Households in Indonesia").
- **Deontic (D):** The extent to which the action is a duty (e.g., "must").
- **Aim (I):** The action being performed (e.g., "reinforce their foundations bi-yearly").
- **Condition (C):** The circumstances under which the action occurs (e.g., "if they live with more than four people").
- **Or Else (O):** The consequence of not acting (e.g., "they will not receive insurance subsidy").

Combining these elements produces rules that can be analysed to understand formal and informal institutions. For instance, the following is a formal rule: "Households in Indonesia (A) must (D) reinforce their foundations bi-yearly (I) if they live with more than four people (C) or else they will not receive insurance subsidy (O)." Formal rules can be found in laws and policy documents, more often than not, they follow this structure (Frantz, C. et al. 2021).

Informal rules describe behaviour and therefore are not enforced by an "Or Else" component. There are two types of informal rules, norms and shared strategies. Norms describe behaviour perceived as a duty or expectation from peers within a demographic. Shared strategies are similar to norms, although they do not have a deontic component, they describe habitual behaviour that has emerged and is not necessarily due to social pressures (Frantz, C. et al. 2021).

The IG offers a structured approach to organising institutional content based on common features such as actions, actors, permissions, and constraints. Analysts can systematically analyse how institutions shape behaviour by parsing institutional directives according to IG syntax.

The literature that will be reviewed in Chapter 4 either uses survey results for non-IG related analyses or performs qualitative IA on interviews and desk research. Commonly, survey results help define relationships between external factors and behaviour, agent-based models, or elements of a social network (Breza, E. et al. 2014). Quantitatively analysing survey data to extract institutions is a novel approach. For this approach, the text of the correlated questions and answers should be converted into ADICO components and the relationships in the survey responses should be identified. However, before this, key definitions in surveys should be outlined.

3.2. Survey definitions

Survey questions provide a wealth of information that can be linked to IG components to extract meaningful insights into social behaviour and institutional rules. To effectively utilise survey data in this context, it is crucial to understand the different types of surveys and survey questions. This section will explore what information survey questions provide and how they can be linked to IG components. There are many types of survey data and ways it can be used for analysis, each offering unique insights into demographic behaviours and trends.

Cross-sectional data is collected at a single point in time, providing a snapshot of a population or phenomenon. This data is useful for identifying prevalence and correlations at a specific moment. For instance, a cross-sectional study might measure disease prevalence at a specific time, whereas a longitudinal study tracks disease development in the same group over several years (Longford, N. T. 2008). Longitudinal data involves repeated observations over time, allowing for the analysis of changes and trends.

Panel data, a type of longitudinal data, contains observations of different cross-sections over time, such as countries, firms, or individuals. This allows for examining changes within the same subjects over time. Non-panel data typically consists of cross-sectional or time series data alone, without tracking the same individuals or entities over time (Grill, C. 2017).

Survey questions come in various forms, each providing different kinds of data that can be used to extract meaningful insights into social behaviour and institutional rules. Regarding the phrasing and ordering of survey questions, it is important to recognise the research available in sociology and psychology (Robinson, S. B., et al. 2018). This research offers valuable insights into framing questions objectively to avoid bias and improve the reliability of responses.

Dichotomous questions offer binary responses, providing clear and straightforward data that is easy to analyse (Robinson, S. B. et al. 2018). They are useful for determining whether certain IG components apply. For example, a question like "Do you intend to purchase flood insurance in the next 12 months?" can indicate the presence of an Aim component. These questions are ideal for rapid understanding and can be easily integrated into decision trees for binary outcomes.

Categorical questions offer multiple predefined response options, allowing for a more nuanced analysis of behaviours, conditions, and attributes (Robinson, S. B. et al. 2018). These questions can help determine various IG components, such as Conditions or

Attributes. Preprocessing categorical data often involves encoding the categories into numerical representations, such as one-hot encoding, to facilitate statistical analysis and language processing.

Ordinal questions, including Likert-scale questions, have a natural order or hierarchy but inconsistent intervals between options. These questions are useful for identifying Conditions or other IG components that rely on progression or ranking. Likert scales measure the intensity or frequency of responses, providing a nuanced understanding of the data (Robinson, S. B. et al. 2018). Preprocessing involves converting responses into numerical values while retaining their natural order.

Open-ended questions allow respondents to provide responses in their own words (Robinson, S. B. et al. 2018), offering deeper insights into norms, values, and shared strategies. These questions require natural language processing techniques for meaningful analysis. This involves tasks such as tokenization, lemmatization, and stop-word removal to extract meaningful information from the text (Khurana, D. et al., 2023). The text must then be analysed for recurring themes, keywords, or phrases that can be linked to IG components.

Survey questions provide critical data that can be structured into IG components, helping to analyse and understand social behaviour and institutional rules. Linking different types of survey questions to IG components is a key part of the standardised method for extracting informal rules. This method will enhance the accuracy and applicability of IA. In conclusion, understanding the types of survey questions and the information they provide is crucial for linking survey data to IG components. The next section will explain decision trees and visualisation techniques to further enhance the analysis of survey data for IA.

3.3. Decision trees

Decision trees are a powerful tool for analysing survey data to identify informal rules and their underlying causes. This section will explain how decision trees work and their potential application in linking survey results to IG components. To understand the application of decision trees in this context, several key variables and concepts need to be defined (James, G. et al., 2023):

- **Node:** A point in the decision tree where the data is split based on a feature (condition). Each node represents a decision point that helps classify the data into different outcomes.
- **Leaf Node:** The endpoint of a branch in the tree, representing a classification or outcome (behaviour) based on the conditions defined in the preceding nodes.
- **Feature:** A condition used to split the data at each node. Features are selected based on their ability to reduce entropy and create homogeneous subsets.
- **Threshold:** A value used to split data at a node. The threshold determines which subset of data moves to the left or right branch of the tree.
- **Class Index:** An identifier for the outcome class at each leaf node, indicating the predominant classification (behaviour) within that subset of data.
- **Entropy:** A measure of impurity or disorder within a sample. Lower entropy indicates a more homogeneous sample, which is desirable for creating clear decision rules.

- **Node Purity:** Similar to entropy, a measure of how homogeneous the data is within a node. Higher node purity indicates that the data within the node predominantly belongs to a single class.

Entropy is calculated by the formula:

$$Entropy = - \sum_{k=1}^K \hat{p}_{mk} \log (\hat{p}_{mk}). \quad Eq.1$$

Where $\hat{p}_{mk}$ represents the proportion of observations in the sample (m) that are from the same class (k) (James, G. et al., 2023). This measure quantifies the level of disorder or impurity within a sample. When all $\hat{p}_{mk}$ values are near zero or one, the entropy will be close to zero, indicating a highly homogeneous sample. Conversely, a higher entropy value suggests a more mixed or impure sample. In the context of decision trees, lower entropy is desirable as it signifies clearer decision rules and more distinct classifications (James, G. et al., 2023).

Decision trees are a type of supervised learning algorithm used for classification tasks. They work by partitioning the data into subsets based on the value of input features, creating a tree-like model of decisions. Each node in the tree represents a feature (condition), each branch represents a decision rule (filtered on the condition), and each leaf node represents an outcome (aim behaviour). This method is useful for handling complex, nonlinear relationships in the data without making strong assumptions about data distribution (James, G. et al., 2023). This process is repeated recursively for each subset of data, creating a tree structure with nodes and branches representing the decision-making process. The goal is to maximise node purity, where each leaf node ideally contains data points belonging to a single class (James, G. et al., 2023). If a leaf has significantly low entropy with enough responses, a relationship can be identified between the final class (aim), and features (conditions) used to reach the leaf.

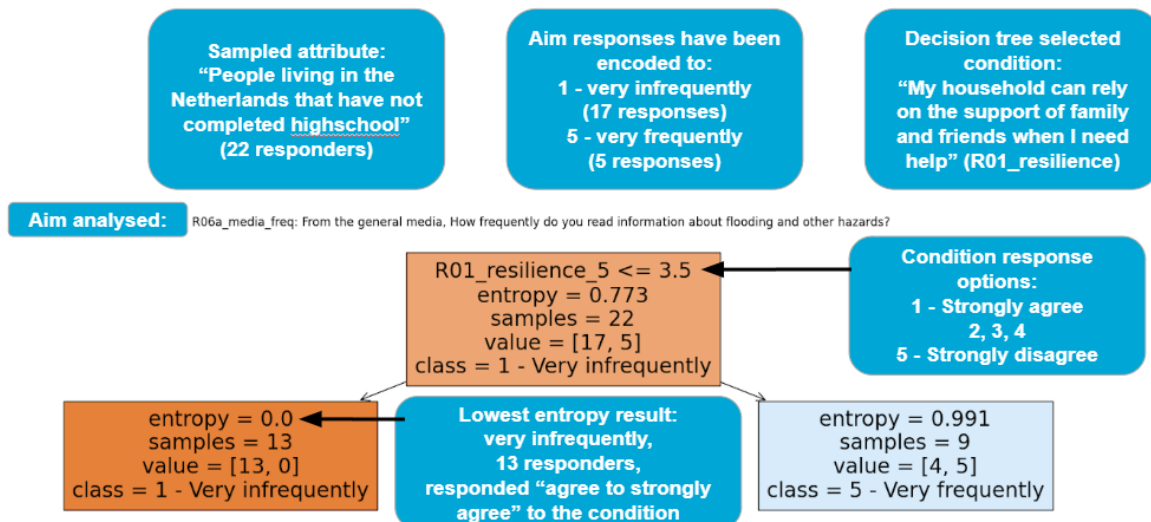


Figure 3.2: Example decision tree with annotations for finding relationships survey results.

Figure 3.2 illustrates a decision tree applied to the SCALAR survey data, demonstrating how responses are categorised to classify behaviours or attributes. This particular tree focuses on respondents who have not completed high school. The root node, represented by the

light-orange box in the centre, contains 22 samples, indicating the total number of respondents with this attribute. The analysis aims to determine how often respondents read information about flooding and other hazards through general media (R06a_media_freq). At the root node, 17 people indicated infrequent engagement, while 5 reported frequent engagement. The decision tree algorithm selects the question that best splits the data to create the most consistent responses to the aim. In this case, it chose a question asking whether respondents believe they can rely on family and friends for support when needed (R01_resilience_5). The left leaf node represents those who believe they can rely on such support, while the right node represents those who don't. The left leaf node, coloured dark orange, indicates a consistent response to the aim question: all 13 respondents in this sub-sample reported infrequently reading information from general media. This uniformity corresponds to an entropy of 0, potentially signifying an informal rule. The right leaf node, coloured light blue, shows a majority of respondents (5 out of 9) frequently reading information from general media. This split corresponds to an entropy value of 0.991. For context, if responses in a sub-sample were equally divided regarding the aim question, the entropy value would be 1. This decision tree visualisation helps identify patterns and potential informal rules within the survey data, providing valuable insights for institutional analysis.

One of the significant benefits of classification decision trees is their high level of interpretability and visualisation. Decision trees are easy to understand and explain, making them suitable for survey data analysis, where the goal is to classify behaviour and identify underlying causes. Decision trees are also robust to outliers, since these trees make decisions based on feature splits rather than fitting a specific function to the data, they are less sensitive to extreme values. This characteristic contributes to more reliable results when analysing survey datasets with potential anomalies (James, G. et al., 2023).

A consideration that needs to be made for extracting data from decision trees is the desired values for sample size and entropy at a node to indicate that the relationship between the variables is strong enough. Even though a node has an entropy value of 0, if only 12 people out of 1000 responders meet the conditions, there might not be enough evidence to point to a pattern of behaviour. In designing the method of this thesis, it will be necessary to test different threshold values to ensure the output is informative and correct. A potential reference value could be from Schelling's model of segregation whereby individuals with a small amount of shared behaviour or preferences can form segregated groups. The threshold in Schelling's model was 33% of an individual's neighbourhood or contacts (Schelling, T. C. 1971).

3.4. Outputs and visualisations

This section explores the potential outputs of the survey data analysis and how they can be effectively visualised to facilitate understanding and application in IA. By converting survey data into IG components, the findings can be presented in a structured manner that highlights key insights and relationships.

Tabular data is a fundamental way to present survey results linked to IG components. A well-structured table can display the relationships between attributes, aims, and conditions derived from survey responses. This allows for a straightforward interpretation of the data, making it easier to identify and analyse informal rules.

The format of Table 2.1 shows a basic way to represent informal rules as IG components extracted from survey data, each of the columns is explained below:

- **Attribute:** Identifies the specific group or entity (e.g., people living in the Netherlands who have not completed high school).
- **Aim:** Describes the intended action (e.g., coordinate with the neighbours in case they are not home when a flood occurs).
- **Condition:** Specifies the circumstances under which the action should take place (e.g. if their household can rely on the support from their government when they need help).
- **Percentage of attribute group in matching statement:** Indicates the proportion of the attribute demographic that matches the aim and condition outcome.

Table 2.1: Example simple table containing output data of informal rules from survey results

Attribute	Aim	Condition	Percentage of attribute group in matching statement
People living in the Netherlands who have not completed high school	coordinate with the neighbours in case they are not home when a flood occurs	if their household can rely on the support from their government when they need help	75%
“““	“““	“““	“““

This basic form of output values is an example and will be improved to become more informative and compatible with IA. A table not only displays the IG components but can provide relevant values and statistics, which are essential for understanding the prevalence and impact of informal rules within the population. While tables offer a clear and detailed view of the data, visualisations can provide more intuitive insights, making it easier to communicate findings to a broader audience.

Bar plots are highly effective for visualising the distribution of responses across different categories. They can illustrate the proportion of respondents who perform specific aims given that they have experienced related conditions. Parallel set diagrams are useful for showing the relationships between multiple categorical variables, revealing how different attributes, aims, and conditions intersect and flow between them. This method provides a clear visualisation of complex data relationships. Sankey diagrams excel at illustrating data flow and highlighting the relative importance of various paths, showing how survey responses split and merge at different decision points, thus emphasising the impact of specific conditions on outcomes (Wilke, C. 2019).

When comparing Sankey and parallel set diagrams, both are effective for visualising flows and relationships between categories, but they differ in their focus and structure. Sankey diagrams emphasise the magnitude of flows between nodes, with the width of the links proportional to the quantity they represent, making them ideal for showing quantitative changes across stages. Parallel set diagrams, on the other hand, focus more on showing the distribution and relationships between multiple categorical variables simultaneously, with equal-width connections that highlight the proportional relationships between categories.

While Sankey diagrams are better suited for processes with a clear direction or flow, parallel set diagrams excel at revealing complex intersections and patterns across multiple variables without implying a specific flow direction (Ribecca, S. 2021).

When visualising proportions described by more than two categorical variables, parallel set diagrams are a preferable choice. In a parallel sets plot, the dataset is broken down by each individual categorical variable, with shaded bands showing how the subgroups relate to each other. For example, a study might break down survey responses by demographic groups, adaptation measures, and perceived barriers, showing the flow of respondents through these categories (Wilke, C. 2019).

In conclusion, this chapter has laid the foundational theoretical framework necessary for this research. By examining IG and its components, a structured approach to analysing informal rules within social systems has been established. Survey definitions and question types were explored to understand how they can be linked to IG components, providing a method for extracting meaningful insights from survey data. The utility of decision trees was discussed, highlighting their ability to classify behaviours and identify underlying causes in survey responses. Various output formats and visualisation techniques, such as tabular data, bar plots, and parallel set diagrams were reviewed to ensure the results are effectively communicated. These theoretical tools will guide the development of a standardised method for extracting and analysing informal rules, enhancing the accuracy and applicability of IA.

4. Literature Review

This chapter conducts an extensive review of existing literature relevant to the research questions, positioning the research context and identifying key theoretical frameworks and concepts. This exploration provides a deep understanding of the landscape surrounding IA methodologies, laying the groundwork for the research methods and development of deliverables.

The literature review process begins with the master theses of Wagenblast (2022) and Lechner (2022) as foundational references. These works analyse the SCALAR survey dataset on flood and climate adaptation, which serves as one of the reference surveys for the deliverables designed in this thesis. To expand the literature base, a systematic search is conducted using databases such as Scopus and Google Scholar, employing specific keywords listed in the abstract of this thesis. Figure 4.1 illustrates the literature collection steps, presented in a left-to-right progression. The initial stages involve discovering literature through recommendations and exploring papers connected to these recommendations, complemented by targeted database searches. This comprehensive approach yields a preliminary list of approximately 90 papers, as depicted in step 3. Subsequently, these papers undergo a rigorous filtering process based on recency and relevance to the subject matter. The final selection, shown in step 5, comprises 20 papers that form the core of the literature review.

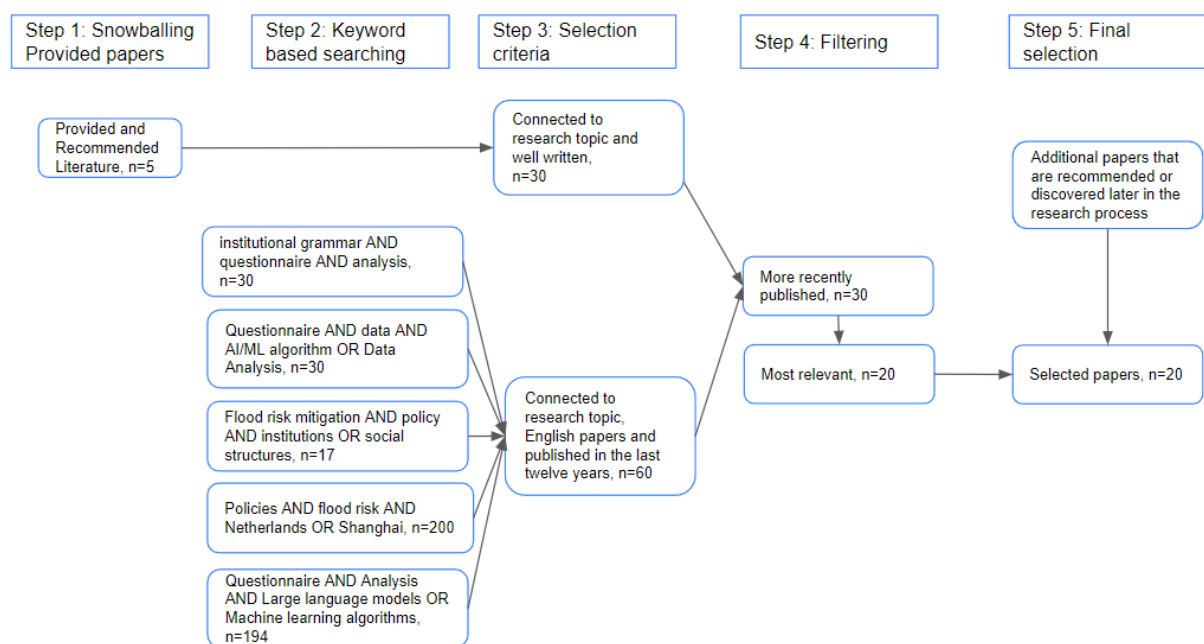


Figure 4.1: Visualisation of the process of scoping down literature towards the final selection for Master Thesis Identifying Institutions and Social Behaviour

It's important to note that this selection process maintains a degree of flexibility. As the research progresses, the relevance of certain papers may shift, and additional literature may be discovered. This adaptive approach ensures that the literature review remains current and aligned with the evolving focus of the research, allowing for the inclusion of newly identified relevant works or the exclusion of papers that become less pertinent to the study's objectives.

Figure 4.2 illustrates the various components and processes of IA, highlighting the transformation of raw research data through various processing approaches into actionable outputs. An objective of this master thesis is to utilise the concepts of this research scope to fill the research gaps identified in the literature review.

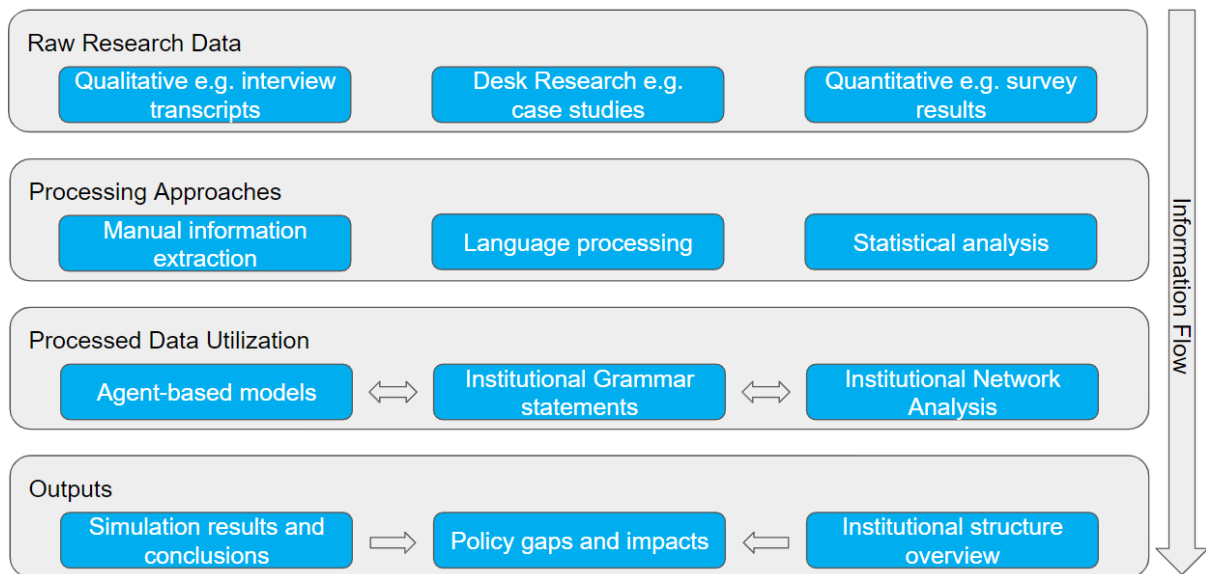


Figure 4.2: Relevant Scope for Master Thesis Identifying Institutions and Social Behaviour

The relevant literature is categorised within the following topics:

- **Challenges in Policy Analysis and Identifying Informal Rules:** Understanding the obstacles policy analysts and institutional analysts face when identifying informal rules.
- **Utilisation of Surveys in Sociological Research:** Investigating how surveys are used in sociological research.
- **Institutional Grammar and Its Applications:** Examining how IG is applied by analysts, the potential outputs, and useful formats for presenting informal rules.
- **Language Processing for Institutional Analysis:** Reviewing relevant literature that utilises language processing techniques.

This structured review aims to uncover gaps and opportunities within the existing methodologies, guiding the development of a robust and standardised approach to extracting informal institutional rules from survey data.

4.1. Challenges in Policy Analysis and Identifying Informal Rules

This section explores the challenges policy analysts and institutional analysts face when identifying informal rules. This exploration is crucial for understanding the current limitations and informing the design requirements for a more effective informal rule extraction method. Institutional analysts currently identify informal rules through qualitative methods, including interview transcripts and desk research. These methods are labour-intensive and subjective, carrying risks of bias and often relying on small sample sizes that do not represent the broader population (Mesdaghi, B. 2020; Verheul, D. 2021; Juarez Pastor, L. 2022). By contrast, formal rules are typically extracted from legal documents through more structured and objective processes (Bognø, J., 2021). The contrast highlights the need for a

standardised methodology to extract informal rules from quantitative data sources like surveys, which can provide more holistic and unbiased insights.

Informal institution discovery requires observations, interviews, simulations, or desk research, and methods such as INA offer approaches to translate qualitative data into institutional diagrams (Ghorbani, A. et al., 2020). Although INA does not include a quantitative element, it offers valuable lessons on language processing for identifying institutional rules. Figure 4.2 shows the steps of the INA method and includes the contribution that this thesis could make. The steps Collection (1), Coding & Clustering (2) and Formalisation (3) should inform the design of the extraction method. Specifically in the way that text is analysed to identify and formulate institutions using ADICO.

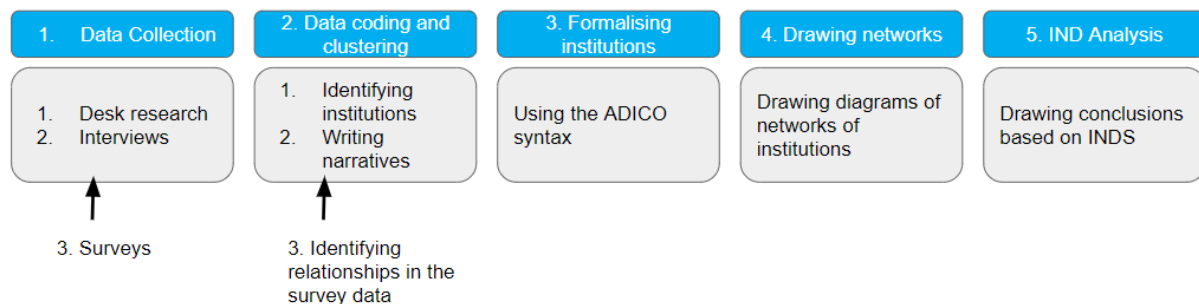


Figure 4.2: The INA Methodology as described by Verheul, D. (2021) with the potential contribution of surveys and data analysis.

The work by Watkins et al. (2016) provides valuable insights into the challenges of extracting informal rules from qualitative data, particularly in the context of ecological restoration decision-making. Their study highlights the difficulty in applying the Institutional Analysis and Development (IAD) framework and ADICO syntax to interview data, as opposed to written policy documents. The authors found that respondents do not articulate rules, norms, and strategies in a straightforward manner, but rather describe their actions and decision-making processes through anecdotes and personal assessments. This necessitated the development of a manual interpretation and extraction process, which involved carefully re-reading interviews and fieldnotes to identify how respondents described typical decision-making within their organisations. Their approach highlights the labour-intensive nature of such work and the potential for subjectivity in interpretation. This reinforces the argument for developing more standardised, quantitative methods for extracting informal rules, which could potentially overcome some of the limitations identified in their study while still capturing the nuanced insights that qualitative data can provide.

Various other studies highlight the importance of improving data collection practices, developing robust formalisation processes for informal rules, and reducing the labour intensity and bias in current methods. For example, Wagenblast (2022) and Lechner (2022) explored private flood adaptation and household roles in reducing coastal flood risk, respectively, and highlighted the limitations of using simplified models and specific data constraints. These studies indicated the need for more representative datasets and refined modelling approaches to capture the complexities of social structures (Wagenblast, T. 2022; Lechner J. 2022).

Abebe et al. (2020) integrated agent-based modelling with flood simulations to study individual adaptation behaviour and flood risk management, demonstrating the significance

of social networks and simple adaptation measures. De Silva et al. (2020) analysed socioeconomic influences on economic loss due to floods, emphasising the need to consider diverse economic groups in disaster risk management. These studies collectively point to the necessity of more sophisticated and holistic approaches to understanding the dynamics of social behaviour and institutional impacts (Abebe, Y. A. et al., 2020; De Silva, M. et al. 2020).

The literature on policy analysis provides further insights into challenges faced in different contexts. Studies examine various aspects of sustainable flood risk management, communication policies, and individual responsibility in flood risk governance (Chan, F. K. S. et al. 2022; Erdlenbruch, K. et al. 2018; Haer, T. et al. 2016; Snel, K. A. W. et al. 2022). These studies highlight the importance of local context consideration, empirical analyses of residents' perceptions, and dedicated studies on implementation barriers to inform effective policy recommendations. A standardised methodology for quantitatively verifying the institutional context could significantly enhance the exploration and understanding of these issues.

In conclusion, the reviewed literature indicates the need for improved data collection practices, the development of a robust informal rule formalisation process, and methods to reduce labour intensity and bias. Key findings from the literature for problem identification are listed in Appendix A. A standardised methodology can address these challenges by providing a consistent, reproducible framework for extracting informal rules from structured data, thereby enhancing the accuracy and applicability of IA.

4.2. Utilisation of Surveys in Sociological Research

This section investigates how surveys are used in relevant research. By examining papers that use surveys, the aim is to identify the synergy between IG and survey data. IG provides a structured approach to organising institutional content, making it a valuable tool for institutional analysts. By parsing institutional directives according to IG syntax, analysts can systematically analyse how institutions shape behaviour. The relevant components of IG are outlined in section 3.1. The structured approach of IG is particularly useful when dealing with informal institutions, such as norms and shared strategies.

To explore the alignment between IG and survey data, literature that uses surveys to analyse socio-behavioural factors is reviewed. Noll et al. (2022), Wagenblast (2022), and Lechner (2022) provide relevant insights through their statistical analysis of the SCALAR dataset. For instance, they examine the effect of socio-behavioural factors on households' climate change adaptation (CCA) intentions, discovering key correlations within the dataset. They also implement sensitivity analyses, such as "odds ratios," to connect components within the survey data. Lechner (2022) uses the odds ratios approach on a subset of the 2020 SCALAR survey data to explore the effects of threat and coping appraisal on adaptation intentions.

The European Social Survey (ESS) is another valuable resource for this research. ESS provides freely accessible data for academics, policymakers, and civil society, allowing for socio-behavioural analysis (European Social Survey, 2024). By reviewing how existing literature utilises ESS data potential strengths and limitations in the context of IA are identified. Chueri et al. (2023) investigate the impact of descriptive and substantive

representation on women's voting patterns for populist radical right parties (PRRPs) using ESS data. The study's strength lies in its use of conditional logit models to analyse voting behaviour in multiparty systems, considering party choice based on available options.

Pavlova et al. (2023) utilise the ESS data to explore conditions that influence actions. They specifically examine the relationship between voluntary participation and well-being across different age groups and European countries using data from Round 6 (2012) of the ESS. The study uses a two-level multiple regression analysis to investigate cross-national variations in voluntary participation and well-being. Despite its strengths, such as mitigating biases through extensive control variables, the study has limitations, including its cross-sectional design and the age of the sample dataset, which may limit the relevance of findings due to significant political and economic changes.

Krefeld-Schwalb et al. (2024) provide valuable insights into extracting patterns of social behaviour from survey data. Their use of regression modelling to explore the relationship between motives and sustainable behaviours should be informative for an approach to identify connections within IG social structures. The study's focus on individual differences emphasises the importance of acknowledging diverse perspectives, which can guide the approach to capturing variations in institutional informal rules. Additionally, the segmentation of individuals based on their motives suggests the utility of clustering techniques for identifying distinct groups in survey data. By taking inspiration from these methods a more informed standardised approach to IG analysis can be designed.

In conclusion, the literature review highlights the potentially significant role of surveys in IA and extracting informal rules. By understanding the applications and limitations of IG and survey data through these studies, this thesis aims to develop a standardised method that leverages their strengths, enhancing the accuracy and applicability in IA.

4.3. Applications of Institutional Grammar

This section examines how IG is applied by institutional analysts, the potential outputs, and useful formats for presenting informal rules. The relevant literature is explored to identify methodologies that align IG with practical survey data. IG is a tool for analysing institutions by decomposing them into base components, providing a framework for understanding social systems and structures (Frantz, C. et al. 2021). The IG Coder is a web application designed to facilitate the interactive encoding of statements within IG syntax (Bognøy, J., 2021). The IG coder addresses the need for a tailored encoding tool for institutional statements, typically found in policies and regulations. The IG Coder's interactive, colour-coded graphic interface visualises the hierarchical structure of statements, making it easier to understand their complexity. While the IG Coder focuses on policy documents, similar approaches can be adapted for survey data.

In the context of IG, processing survey text to identify components such as Attributes, Aims, and Conditions is critical for converting the text into Institutional Behavior Statements. Several approaches are available for this task, each with unique advantages and challenges. Manual processing involves human interpretation and extraction of ADICO components from survey text, offering high clarity but being slow, labour-intensive, and not scalable for a large number of questions. While existing literature employs significant uses of IG to study social

groups and structures, this thesis explores the potential for institutional analysts. Mesdaghi (2020), Verheul (2021), and Juarez Pastor (2022) implement the INA method to extract institutional structures from qualitative data such as interviews or desk research.

Siddiki et al. (2023) explore the utility of IG in analysing social systems by examining institutional dynamics. Their work highlights how IG can be applied to study formal and informal rules captured in public policies and social conventions. They discuss the evolving trends in IG research, emphasising the importance of understanding institutional changes and adaptations over time and demonstrate how IG can be used to assess environmental governance and policy changes, illustrating the versatility of IG in various contexts (Siddiki, S. et al. 2023).

The methodology from Krefeld-Schwalb et al. (2024) provides an excellent example of how motives for sustainable behaviour can be extracted from survey responses. Though they do not explicitly implement IG, there are significant parallels to draw. The visual representation of the relationship between motives and sustainable behaviours is represented in Figure 3.3.

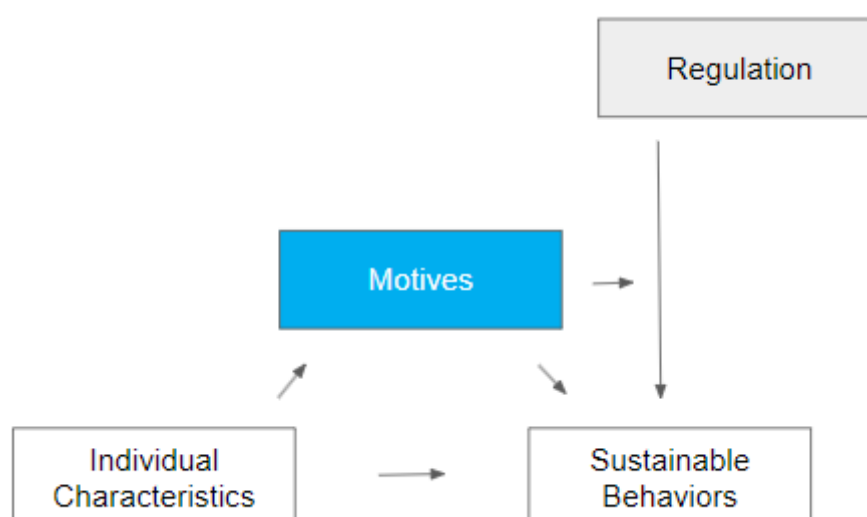


Figure 3.3: Diagram with relationships of motives for sustainable behaviour from (Krefeld-Schwalb et al. 2024)

The process described in the paper for extracting meaningful information from survey responses aligns with the aims of this research, offering inspiration for the design of the extraction method. The diagram in Figure 3.3 from their study can be adapted to represent IG components, substituting motives for conditions and sustainable behaviours for aims.

In conclusion, the reviewed literature shows the potential of IG to structure and analyse informal institutions within survey data. By combining IG components with advanced data analysis techniques, this thesis aims to develop a standardised method for extracting informal rules, enhancing the accuracy and applicability of IA. This approach addresses the limitations of existing methodologies and contributes to a deeper understanding of social behaviours and institutional impacts.

4.4. Language Processing for Institutional Analysis

This section reviews relevant literature that utilises language processing techniques, examining how these techniques could enhance IA and IG. By identifying methods to extract IG from survey data, existing approaches are better understood, establishing a framework for converting survey responses into IG syntax components.

The literature reviewed thus far mainly uses survey results for non-IG related analyses or performs qualitative IA on interviews and desk research. Commonly, survey results help define relationships between external factors and behaviour, agent-based models, or elements of a social network (Breza, E. et al. 2014). For analysing survey data to extract institutions, the text of the correlated questions and answers should be converted into ADICO components. This section outlines key literature on language processing techniques. Creating a list of relevant terms and phrases for each ADICO component and using a dictionary or semantic search for automated text analysis can improve efficiency. This approach is similar to research where NLP was applied to peer-reviewed publications to detect shifts in farmers' climate change adaptation strategies (Gil-Clavel, S. et al. 2023).

The methodology from Krefeld-Schwalb et al. (2024) offers an excellent example of how motives for sustainable behaviour can be extracted from survey responses using advanced NLP techniques. They utilised an LLM to encode behaviours into neural representations, followed by dimension reduction with clustering to identify key motives. This process, although not explicitly implementing IG, aligns with this research's aims, providing inspiration for the extraction method.

When analysing large volumes of text data, machine learning algorithms offer promising approaches to uncover patterns and insights. However, selecting the right algorithm and training data is crucial for reliability. While natural language processing and LLMs are versatile tools, caution must be exercised due to potential biases (Bender, E. M. et al. 2021). Neural networks have transformed natural language processing, enabling more sophisticated applications across various language-related tasks (Khurana, D. et al., 2023). Language models like those available on HUGGINGFACE can be trained specifically for IG tasks. Bi-directional Encoder Representations from Transformers (BERT) marked a significant advancement by allowing simultaneous examination of text in both directions and contextual embedding for each word. However, BERT's limitation is handling long text sequences, requiring them to be divided into shorter sequences (Khurana, D. et al., 2023). BERTopic encoding applies machine learning to extract topics and patterns from text which could be a viable approach with proper model training.

Artificial Intelligence (AI) is capable of producing realistic text, images, other human-like outputs, and is currently transforming various industries. This technology has the potential to improve social science research methodologies, such as survey research, online experiments, automated content analyses, and agent-based models (Bail, C. A. 2024). However, there are significant limitations, including bias in the training data, ethical concerns, replication challenges, environmental impacts, and the proliferation of low-quality research. Addressing these limitations requires the creation of open-source infrastructure for research on human behaviour (Bail, C. A. 2024). This infrastructure is essential not only to

ensure broad access to high-quality research tools but also to deepen the understanding of the social forces that guide human behaviour.

Generative AI models such as OpenAI's GPT, Gemini, and Mixtral rely on crafting prompts and providing contextual information to guide the model's interpretation of input data (Rangapur, A. et al. 2024). General-purpose models like GPT-3.5, while offering broader coverage, may identify elements with limited relevance to the specific context. In contrast, specialised models focus more on details (Hajikhani, A., et al. 2024). An example of LLMs in urologic research is demonstrated by Kaufmann, et al. (2024), who developed and validated a zero-shot learning NLP tool based on OpenAI's GPT-3.5. This tool facilitated data abstraction from unstructured text in electronic health records, significantly reducing the time required for data abstraction compared to human abstractors while maintaining comparable accuracy. This study highlights the potential of LLMs to enhance the efficiency and reliability of data processing in medical research, emphasising the importance of specialised training for precise and unbiased analysis (Kaufmann, B. et al., 2024).

This research explores language processing techniques alongside core statistical analysis methods to connect ADICO elements within survey data. By prioritising the most accessible, consistent, and accurate approaches, the goal is to develop a standardised method for extracting informal rules, enhancing the accuracy and applicability of IA. Once informal rules have been extracted, it is crucial to consider what data and information are needed for institutional analysts to ensure the results are usable, consistent, and well-communicated.

Chapter 3 established the necessary theoretical background for this research, including the exploration of theoretical tools such as IG and its components, methodologies for linking survey data to IG concepts, and relevant statistical and language processing techniques. This theoretical foundation provides a robust framework for the development and application of the standardised informal rule extraction method. The following chapter introduces the novel approach for extracting informal rules from survey data. This method, termed the Survey Informal Rules Extraction (SIRE) method, operationalizes the theoretical concepts discussed and addresses the research objectives outlined in this study. By providing a step-by-step process for researchers, SIRE aims to bridge the gap between theoretical understanding and practical application in the field of institutional analysis.

5. Proposing the SIRE Method: Extracting Informal Rules from Survey Data

This chapter describes each stage of the proposed Survey Informal Rules Extraction (SIRE) method, outlining what the researcher is expected to do and the anticipated results. The process involves multiple steps, starting from the selection of relevant survey questions to the formation of ADICO statements. The methodology is structured to facilitate the extraction of informal rules-in-use from survey data for IA. A diagram showing the SIRE method steps can be found in Appendix B. The researcher is expected to use Python 3.11.5 for each of these steps (Python. 2024). The standardised Python code to perform these steps, along with the packages to install using pip, can be found in Appendix C. The code has also been prepared and made available in Github (Gittins, G. 2024). Figure 5.1 shows the key steps in the SIRE method which will be explained in the upcoming sections.

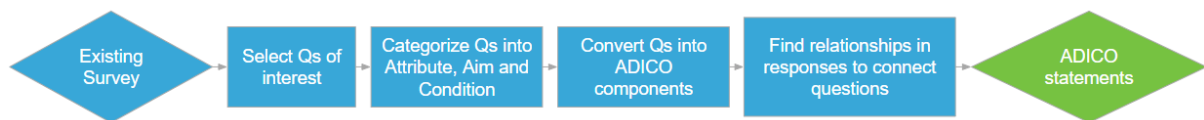


Figure 5.1. Simplified overview of the steps of the SIRE method.

Appendix D contains example data and outputs that should be produced in key steps of the SIRE method. The data used for the figures and examples in this chapter and in Appendix D is from a mock survey. The mock survey only serves as a standardised example that is compatible with the SIRE method, no conclusions can be drawn from the results.

5.1. Survey Data Examination

To begin implementing the Survey Informal Rules Extraction (SIRE) method, it is important to explore and prepare the survey data. This data should include questions relevant to the research context, structured in two primary datasets: a question overview and responses. First make sure to install the necessary python packages, pandas and numpy using “pip install”.

Question Overview Data: This dataset should include the following columns:

- **Question_ID:** A unique identifier for each question.
- **Question_Text:** The text of the question.
- **Response_Option:** The text of each response option.
- **Response_Value:** Numerical mapping of response options.
- **ADICO_Category (Optional):** The ADICO component (Attribute, Aim, Condition, Deontic) identified by the question.

Including the ADICO_Category column is optional because it may not be necessary for all analyses. However, when available, it helps categorise questions into specific components, facilitating more targeted analysis in later steps. For example, it can be used to inform a generative AI model how to codify the text as a respective ADICO component.

Response Data: The response data should consist of rows representing individual responses and columns representing question identifications. Ideally, the response values should be numerical, allowing them to match with the text values in the survey overview data.

Steps to Examine Survey Data:

1. **Access Survey Data:** Ensure the availability of both the survey overview and response datasets. For instance, load the CSV files containing the survey data as shown in Appendix C1.
2. **Review Survey Overview Data:** Ensure the question overview data includes all necessary columns (Question_ID, Question_Text, Response_Option, Response_Value, and optionally ADICO_Category).
3. **Check Response Data:** Verify that the response data contains numerical values for each question, enabling a match with the text values in the survey overview data.
4. **Generate Unique IDs for Response Options:** Create a unique identifier for each response option, facilitating easy matching between responses and question texts. For example, question 1's response value 1 could be identified by the code "Q1_1".

5.2. Selecting and Categorising Questions

To effectively implement the SIRE method, the next step involves selecting and categorising questions from the survey data that are relevant to the research objectives. This process ensures that the analysis captures the necessary representation of demographics, actions, opinions, and external factors.

Initially, samples of the response data should be created based on specific responses to attribute questions. Attributes, while not always necessary, can help focus the analysis on particular subsets of the data, such as specific demographic groups. For example, responses might be filtered based on demographics such as age, location, or educational level to observe trends within these groups.

Next, select questions from the survey data that are relevant to the research objectives. The selection should ensure a holistic representation of demographics, actions, opinions, and external factors. Prioritise dichotomous and ordinal questions due to their simplicity in coding and analysis.

Steps for Selecting Questions:

1. **Demographic Questions:** If the focus is on studying a specific demographic, include attribute-categorised questions that capture relevant demographic information.
2. **Action and Opinion Questions:** Identify questions that reflect actions performed or intended, opinions held, or changes made by respondents.
3. **External Factors:** Include questions that capture external factors influencing behaviour, such as environmental conditions or societal influences.
4. **Assign ADICO Categories:** If the ADICO_Category values have not been assigned yet, categorise the selected questions into ADICO components:
 - **Attributes:** Questions that distinguish relevant demographics.

- **Aims:** Questions about actions performed or intended, opinions held, or changes made.
- **Conditions:** Questions about external factors influencing behaviour.

Refer to Appendix C2 for additional details and examples of the code implementation.

5.3. Data Preparation and Preprocessing

The data preparation and preprocessing stage is crucial for ensuring the quality and reliability of the subsequent analysis. During this phase, the survey data is prepared for both parallel set diagrams and decision tree analysis. This process involves several key steps: data cleaning, response encoding, and formatting the data for optimal analysis.

The first step in data preparation is a thorough cleaning process. This involves carefully examining the dataset to identify and exclude any response options that do not contribute to meaningful outcomes. Responses such as "n/a" or "Don't Know" often fall into this category, particularly if the remaining response options are ordered. The inclusion of these non-contributory responses can significantly skew analysis results, potentially leading to misinterpretations. In cases where data points are missing, researchers may need to consider employing "missing data imputation" techniques. These methods involve estimating missing values based on other available data, helping to maintain the integrity and completeness of the dataset. Following the cleaning process, the next critical step is the manual encoding of aim and condition responses. This encoding process typically involves converting textual or categorical responses into binary or ordinal values. This step is crucial for two primary reasons: firstly, it simplifies the data for decision tree analysis, and secondly, it ensures that the responses are interpretable and consistent across the dataset.

To illustrate this encoding process, let's consider a specific example. For a survey question such as "Do you intend to purchase flood insurance in the next 12 months?", responses can be encoded into a binary format. In this case, affirmative responses indicating an intention to purchase would be encoded as "yes" or "intend to", while negative responses would be encoded as "no" or "don't intend to". This binary encoding allows for clear differentiation between the two possible outcomes.

Figure 5.2 provides a visual representation of this encoding process. It demonstrates how survey responses are systematically converted into numerical values for analysis. In this illustration, responses that indicate the implementation of a specific behaviour are encoded as 1 and represented in green. Conversely, responses that do not indicate the implementation of the behaviour are encoded as 2 and represented in red. This colour-coding system not only aids in the visual interpretation of the data but also facilitates the identification of patterns and trends in the responses.

Being an active member in a community group aimed at making the community safer	1	I have already implemented this non-structural measure
	2	I intend to implement this non-structural measure in the next 6 months
	3	I intend to implement this non-structural measure in the next 12 months
	4	I intend to implement this non-structural measure in the next 2 years
	5	I intend to implement this non-structural measure in future, after 2 years
	6	I do not intend to implement this non-structural measure

Figure 5.2. Encoding survey response options, green becomes 1: Has implemented and red becomes 2: Has not implemented

Similarly, select condition questions and encode their responses as binary or ordinal outcomes, depending on the nature of the questions. For example, questions about the frequency of certain behaviours can be encoded to reflect their ordered nature. Properly encoded data can be easily used in parallel set diagrams and other visualisations, aiding in the clear communication of results.

By following these data preparation and preprocessing steps, researchers can ensure that their survey data is optimally structured for analysis. This careful preparation lays the groundwork for more accurate and insightful parallel set diagrams and decision tree analyses, ultimately contributing to a more robust extraction of informal rules from the survey data. Refer to Appendix C3 for further details and examples of the code implementation.

5.4. Rewriting to ADICO Components

In this section, the process of converting survey questions and response options into their ADICO component equivalents is described. The ADICO framework helps in structuring institutional statements by breaking them down into Attributes, Deontics, Aims, Conditions, and Or else. For each selected question and response option, translate the text into its respective ADICO components. Here is an example of how to break down the components:

- **Attribute:** Identifies the specific group or entity involved.
- **Aim:** Describes the intended action or behaviour.
- **Condition:** Specifies the circumstances under which the action occurs.

If there are many questions to process, consider using AI to automate the conversion. Appendix C4 provides an example code using an LLM to perform this task. Currently, Groq can be used to make text generation requests for free (Groq. 2024). Making requests is similar for the different services so with few adjustments to the code, OpenAI or Claude can also be used (Rangapur, A. et al. 2024). Automating this step can save time for large numbers of questions and ensure consistency, although manual conversion can be used to avoid errors.

Example Question, Response and Conversion:

- Attribute Question: "What is your Age?"
 - Response Options: "50 or younger", "Over 50"
 - **Attribute:** People over the age of 50
- Aim Question: "Do you intend to install solar panels?"
 - Response Options: "Yes", "No"

- **Aim:** Intend to install solar panels
- Condition Question: "Will you receive a subsidy if you install solar panels?"
 - Response Options: "Yes", "No"
 - **Condition:** If they receive a subsidy for installing solar panels

By following these steps, researchers can systematically rewrite survey responses into structured ADICO components, facilitating the extraction and analysis of informal rules from survey data. Refer to Appendix C4 for detailed examples and code implementations for this task.

5.5. Creating Parallel Set Diagrams

This section will describe how to visualise the selected survey data for preliminary exploration and analysis using parallel set diagrams. Parallel set diagrams are highly effective for representing the initial distribution of responses and identifying patterns among Attributes, Aims, and Conditions. This visualisation aids in understanding the relationships and flow of responses across different categories, making it easier to analyse and interpret the data.

Steps for Creating Parallel Set Diagrams:

1. **Install Plotly:** Using pip install, install the necessary Plotly package (5.20.0)
2. **Define Custom Colour Mapping:**
 - a. **Assign Colours:** Assign colours to the aim question response categories to help visualise the distribution of the aim responses across each of the condition question response categories.
3. **Adjust Labels:**
 - a. **Modify Aim Labels:** Modify the labels of the aim responses to include counts and percentages to make them more informative.
 - b. **Modify Condition Labels:** Adjust the labels of the condition responses to include counts and percentages to provide additional context.
4. **Create the Parallel Set Diagrams:**
 - a. **Configure Plot Headers:** Define the headers for the diagram based on the selected questions.
 - b. **Plot the Diagram:** Using the response data, create the parallel set plot.
 - c. **Customise Layout:** Update the layout of the figure with a title, font size, and other aesthetic configurations.
 - d. **Display the Plot:** Display the plot to visualise the data.

Figure 5.3 shows a parallel set diagram created from responses to the SCALAR dataset. The title of the diagram indicates the attribute demographic value and the aim. The split on the far-right represents the proportion of responders that indicated intending to perform the aim and not intending to, these are also colour-coded as blue and grey respectively. Moving to the left the responses are split according to different condition questions where different proportions can be observed regarding the aim response. This provides the researcher with an overview of the variation of response to the aim behaviour given their response to the other questions. This helps to contextualise the statements that will be extracted as well as

inform the researcher which factors to pay attention to and identify any necessary further data processing.

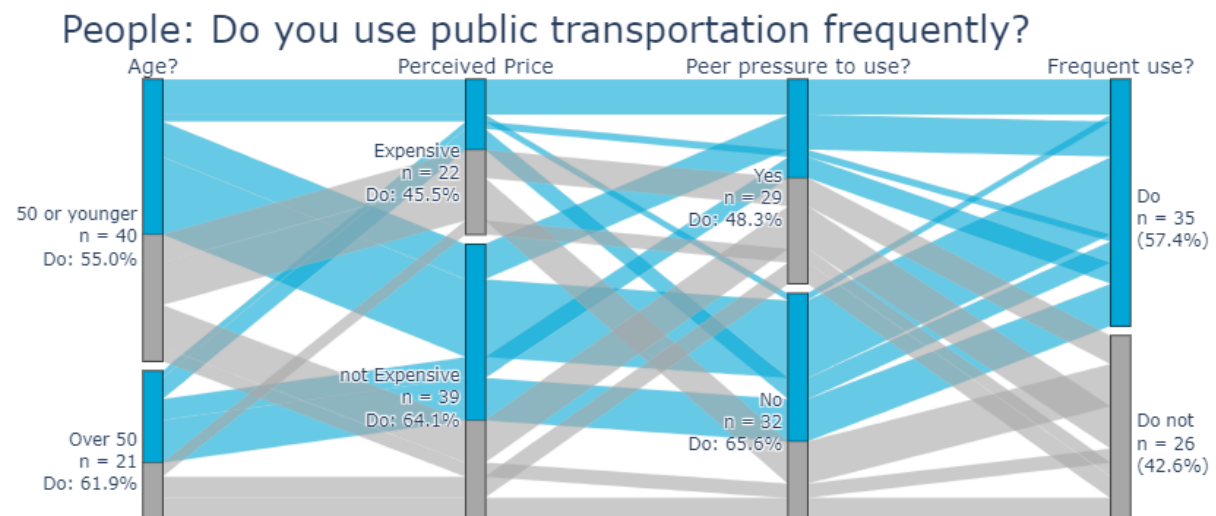


Figure 5.3: Example Parallel Set Diagram produced by using the script in Appendix C5 on the SCALAR dataset

These steps create a parallel set diagram that visually represents the distribution of responses and highlights patterns among Attributes, Aims, and Conditions. This visualisation tool is essential for understanding complex relationships in survey data and facilitating more insightful analysis. For more detailed examples and the full code, refer to Appendix C5.

5.6. Extracting Statements using Decision Trees

With the data prepared and the parallel set diagrams produced and interpreted, each sample is processed using a decision tree algorithm, iterating over the aim questions. The decision tree algorithm works by partitioning the responses based on the value of input features (conditions). Each node in the tree represents a feature, each branch represents a decision rule, and each leaf node represents an outcome. The algorithm selects the condition to split the responses to the aim question and find a leaf with the least entropy, indicating a strong relationship between the condition and the aim. The algorithm practises conditions that result in the most homogeneity in responses to the aim question.

An important consideration when using decision trees is that they will select the optimal condition to split the responses and find a leaf with the least entropy, thereby identifying the local maxima condition(s) related to the aim. To discover the second and third best conditions related to the aim, the process is repeated, removing the previously selected condition from the options each time. This iterative approach ensures a comprehensive analysis of all relevant conditions influencing the aim.

1. **Training the Decision Tree:** The decision tree is trained using survey responses, where features are the condition responses and target outcomes are the aim responses. The tree partitions the data based on these features, creating branches that represent different decision paths.
2. **Node Evaluation:** The tree's nodes are evaluated based on entropy and sample size to determine the purity of the splits. Nodes with lower entropy indicate a more

homogeneous response to the aim question and nodes with large sample sizes indicate a larger number of responders that share the conditions and aim outcome.

3. **Identifying Significant Nodes:** Nodes that satisfy thresholds for entropy and sample size are identified. These nodes represent points in the data where the relationship between attributes, aims, and conditions are sufficient to indicate potential informal rules. These threshold values can be set at the researchers discretion. Suggested values are a sample size of at least 33% of the original sample and an entropy of at most 0.75.
4. **Extracting IG Components:** For each significant node, the relevant IG components (Attribute, Aim, Condition) are extracted along with statistical data such as sample counts after condition splits and percentage of responders matching the aim response outcome.
5. **Generating Statements:** The extracted information is used to generate structured statements that capture the identified informal rules. These statements are compiled into a table, providing a clear and organised view of the data.

The extraction of informal rules can be fine-tuned through iterative adjustments of the sample size and entropy requirements. This iterative approach allows researchers to balance the robustness of the identified relationships with the breadth of extracted statements. If the initial analysis yields no statements, it may indicate that the set requirements were overly stringent or that the decision tree algorithm failed to detect significant relationships within the data. In such cases, researchers can adjust the parameters to better suit their dataset and research objectives.

The sample size requirement plays a crucial role in determining the strength of the extracted relationships. A larger sample size relative to the initial number of respondents implies a stronger, more representative relationship. By increasing the sample size threshold, researchers ensure that statements are only extracted when the node contains a sufficiently large proportion of respondents. This approach enhances the reliability of the extracted rules but may result in fewer overall statements. Similarly, the entropy requirement can be adjusted to control the homogeneity of responses at each node. By decreasing the required entropy value, more statements can be extracted, albeit with a potential increase in the proportion of respondents at each node who don't match the aim behaviour. This trade-off between quantity and quality of extracted statements allows researchers to tailor the analysis to their specific needs.

These adjustments should be made thoughtfully, considering the overall research objectives and the nature of the dataset. The goal is to strike a balance between extracting meaningful, reliable informal rules and capturing a comprehensive view of the relationships present in the data. By carefully calibrating these parameters, researchers can optimise the SIRE method's effectiveness in uncovering informal rules from survey data.

Figure 5.4 shows a decision tree produced in Python using the SIRE method. Each box represents a node which can be interpreted as a potential statement, they contain information about the condition used to make further splits, the entropy at this node, the number of responders that have reached this node, the distribution of responses to the aim at this node and the majority response to the aim at this node. The further down the tree, the more condition questions are used to split the responses, leaf nodes at the bottom do not

have a condition used to make further splits. The lowest entropy value in Figure 5.4 is 0.773 which is indicated by the box with the darkest shade. The boxes are shaded blue if the majority of responders at that node indicate not performing the aim and are shaded orange if the majority do perform the aim.

Q4: Do you use public transportation frequently?

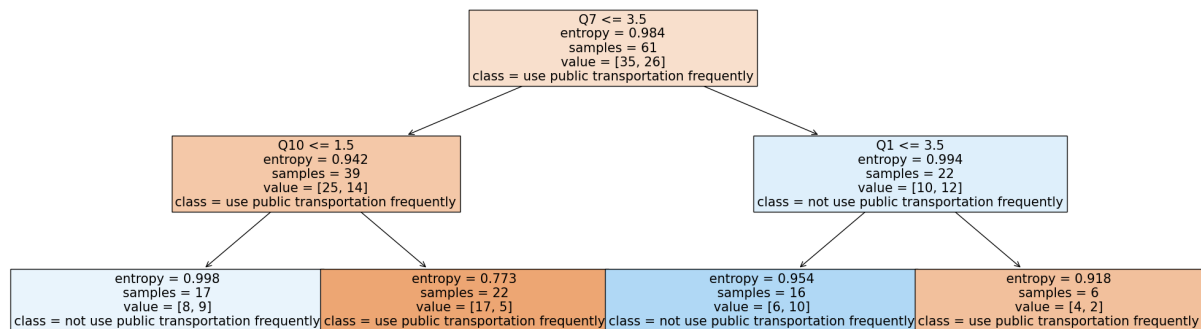


Figure 5.4: Example decision tree visualisation for extracting informal rules from results of a mock survey results

Below is a list of the current output data columns where each row represents a node with a sample size and homogeneity meeting the predetermined threshold requirements and therefore describes an informal rule extracted from the decision tree:

- **Attribute:** Attribute text in IG Syntax
- **Aim:** Aim text in IG form
- **Total_count:** Total number of people in the demographic
- **Aim_%_True:** percentage of people in the total count that perform the aim
- **Condition1:** Text of the first condition used to increase proportion of people performing aim
- **Condition1_count:** number of people in the demographic that satisfy the first condition
- **Condition1_Aim_%:** percentage of people in the condition1 count that perform the aim
- **Condition2:** Text of the second condition used to increase proportion of people performing aim
- **Condition2_count:** number of people in the demographic that satisfy the first and second condition
- **Condition2_Aim_%:**percentage of people in the condition2 count that perform the aim

See Appendix D6 for an example of a table with this extracted data.

5.7. Visualise Extracted Statements

To effectively communicate the significance of the informal rules identified through the decision tree analysis, it is essential to visualise the extracted statements and their proportions. This section provides a detailed explanation of how to perform this visualisation using Sankey diagrams, which are well-suited for this purpose due to their ability to illustrate flows and relationships between categories. The data is leveraged from the decision tree to

create visual representations that highlight the relationships between attributes, conditions, and aims.

Before creating the visualisations, ensure that the data extracted from the decision tree is well-structured and contains the necessary columns listed in the previous section. The primary goal of the visualisation is to display the proportions of respondents meeting the selected conditions and their corresponding aims. Sankey diagrams are used to illustrate these relationships clearly.

1. Initialise the Plot: For each statement in the data, initialise a Sankey diagram using the Plotly plotting library. Set the figure size to ensure the plot is readable.
2. Define Nodes and Links:
 - Create nodes representing the total count of respondents, each condition, and the final outcomes.
 - Define links between nodes to show the flow of respondents through conditions to outcomes.
 - Calculate the values for each link based on the proportions of respondents meeting each condition and aim.
3. Set Colors and Labels:
 - Use distinct colours to differentiate between flows representing those who meet the aim, those who do not, and those excluded by conditions.
 - Label each node with appropriate descriptions and counts.
4. Add Annotations and Legends:
 - Include annotations below the diagram to describe the attribute, conditions, and final outcomes.
 - Add a legend to explain the colour coding used in the plot.
5. Set Titles and Layout:
 - Title each plot with the specific aim it represents.
 - Adjust the layout for optimal readability, including font sizes and diagram dimensions.

Figure 5.5 shows a Sankey diagram used to visualise an informal rule extracted from a survey. The flows represent the number of respondents moving through each condition, with colour-coding to indicate whether they meet the aim, do not meet the aim, or are excluded by conditions. The plot title represents the aim observed, and annotations provide context for each stage of the flow.

Sankey plot representation of an IG statement extracted by decision tree algorithm:
use public transportation frequently

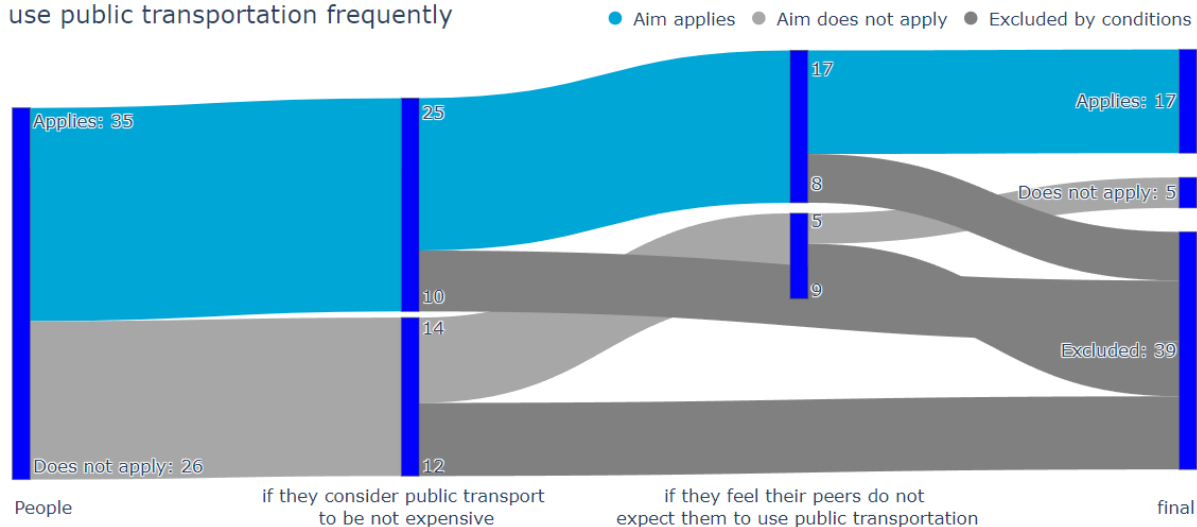


Figure 5.5: Example Sankey diagram visualisation of an informal rule extracted from a mock survey.

The Sankey diagram in Figure 5.5 should be read from left to right, illustrating the flow of survey respondents through various filtering conditions. It begins with the entire surveyed population on the left, divided into two groups: 35 people for whom the aim applies and 26 for whom it doesn't. Moving rightward, the population is filtered based on specific conditions such as low price perception and lack of peer pressure, as identified by the decision tree algorithm. The width of each flow represents the number of people in that category, with colours distinguishing the aim response or condition exclusion. At each stage, the diagram shows how the proportion of respondents for whom the aim applies changes, with the goal of making this response more uniform. The final column on the right displays the resulting distribution after all conditions have been applied: 17 people to whom the aim applies and 5 to whom it does not. This allows for a comparison between the end state (17:5) and the initial distribution (35:26), demonstrating how the selected conditions influence the prevalence of the aim. This visual representation enables a grasp on how the selected conditions might influence behaviour related to the aim, potentially revealing factors that could make the aim more common or relevant among the target demographic. Throughout the diagram, annotations and labels provide context and specific numbers, guiding the reader's understanding of the data flow and its implications.

This approach provides a clear and intuitive visualisation of how respondents flow through the conditions and how this affects the proportion meeting the aim, offering valuable insights into the structure and impact of the extracted informal rules. These steps guide the creation of clear and informative representation of the informal rules extracted from the decision tree analysis, facilitating a better understanding of the relationships between attributes, conditions, and aims within the data. Appendix C7 contains the full Python script that creates these visualisations, providing a practical example of how to visualise the extracted statements.

6. Validating the SIRE Method

In this chapter the SIRE method will be validated through tests on the SCALAR and ESS datasets. First, for each survey a case study will be selected, key results from the case will be identified and compared with the results that the SIRE method produces. Validating the method involves demonstrating that it can produce similar and logical conclusions to those found in the literature, ensuring the reliability and robustness of the approach. Secondly, existing or proposed policies linked to the survey questions will be identified. By comparing the extracted informal rules with these policies, the effectiveness of the policies in influencing the intended behaviours can be evaluated. This comparison will help illustrate how the method can be used to assess and refine policy interventions.

6.1. Validation 1: Comparison with SCALAR Literature Results

The case study selected for the first validation is a study on household climate change adaptation (Noll, B. et al. 2022). They calculated the odds-ratio relationships of survey responders' intentions to implement climate change adaptation measures with factors such as the responders' threat and coping appraisal, background and climate-related beliefs. They present their results in plots such as Figure 6.1 below. For the validation, the SIRE method will make use of the data from Responders in the Netherlands.

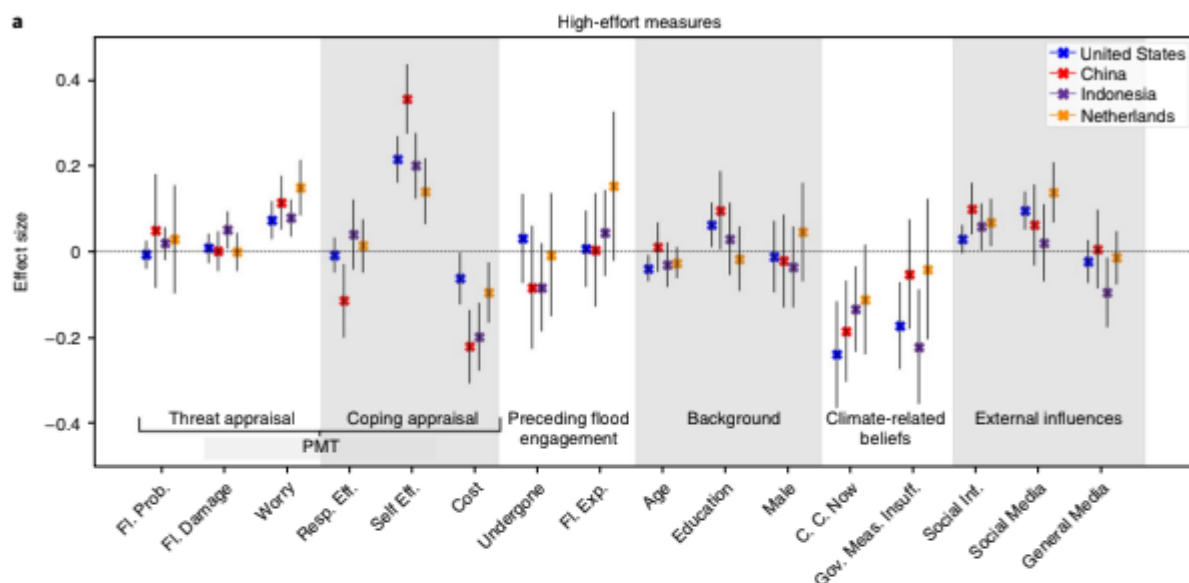


Figure. 6.1. Effect Distributions for Factors Influencing Households' Intentions to Adapt to Flooding (Noll, B. et al. 2022)

Figure 6.1 displays the mean effects for factors influencing households' intentions to adapt to flooding, categorised by high-effort and low-effort measures. The results are derived from Bayesian beta regression models run separately for four countries: the United States, China, Indonesia, and the Netherlands.

The SIRE method will be validated by selecting 3-4 high-effort and 3-4 low-effort measures and seeing if similar relationships emerge. The scope of tested conditions will be reduced to

self efficacy, perceived cost, worry, and age. Based on the results (Noll, B. et al. 2022), the patterns that are expected are as follows:

- **Self-efficacy** should increase the responders' intention to implement high-effort measures and slightly decrease their intention to implement low-effort measures.
- **Perceived cost** should decrease the responders' intention to implement high-effort measures and slightly increase their intention to implement low-effort measures.
- **Worry** should increase the responders' intention to implement both high and low-effort measures.
- **Age** should have little to no effect on the intention outcomes.

The survey responses will be pre-processed to:

Implementation: *Do you intend to implement the following measures?*

1. I intend to implement this measure.
2. I do not intend to implement this measure.

Self-efficacy: *Do you have the ability to implement the structural measure?*

1. I am not able
2. I am able

Perceived cost: *Do you believe that implementing or paying someone to implement this structural measure would be cheap or expensive?*

1. I believe it would be cheap.
2. I believe it would be expensive.

Worry: *How worried or not are you about the potential impact of flooding on your home?*

1. I am not worried
2. I am worried

Age: *Age?*

1. Younger than 45
2. Older than 45

Then by following the steps of the SIRE method, the following parallel set plots are produced. In Figures 6.2 and 6.3, clear patterns emerge regarding intentions to implement a high-effort measure to raise the ground floor level above the most likely flood level in the Netherlands.

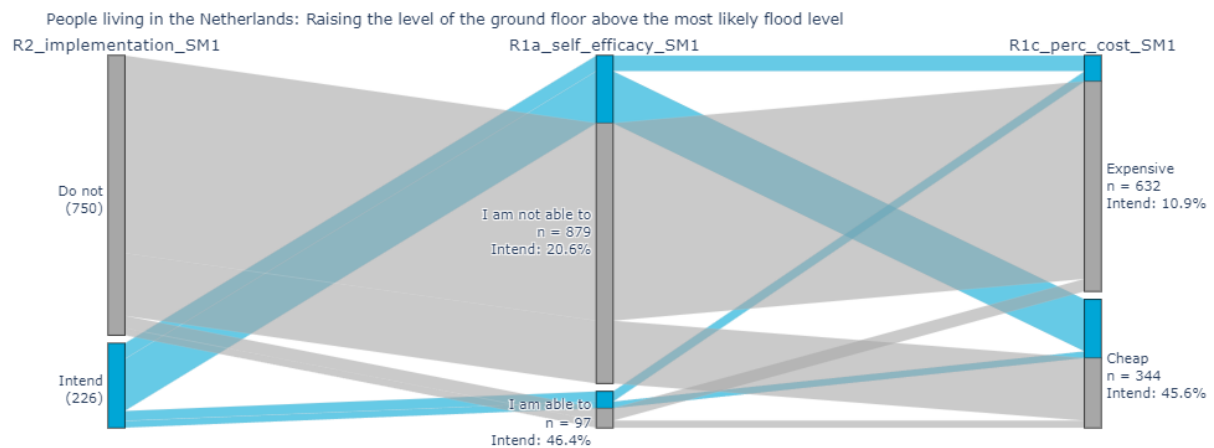


Figure. 6.2. Parallel set diagram for SIRE SCALAR Netherlands analysis - relationships between individuals' intentions to raise the level of the ground floor above the most likely flood level (high-effort), self-efficacy and perceived costs.

Figure 6.2 reveals that a significant majority (76.8%, 750 out of 976) do not intend to implement this measure. The data also shows that 90.1% (879 out of 976) of respondents are not able to implement the measure. However, among those who are able to implement it, the intention rate is much higher at 46.4% (45 out of 97). Similarly, for those who perceive the measure as cheap, the intention rate increases to 45.6% (157 out of 344), compared to only 10.9% (69 out of 632) for those who view it as expensive.

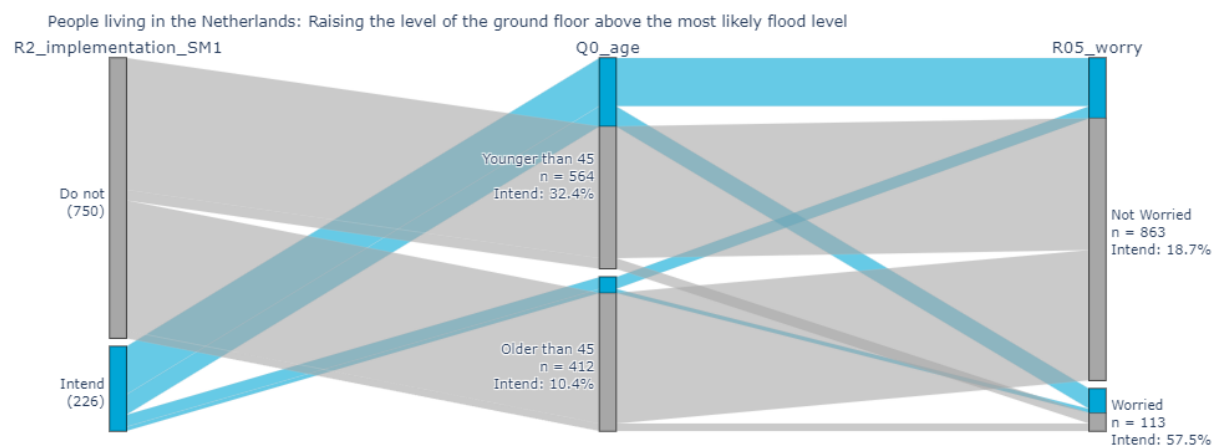


Figure. 6.3. Parallel set diagram for SIRE SCALAR Netherlands analysis - relationships between individuals intentions to raise the level of the ground floor above the most likely flood level (high-effort), age and worry.

Figure 6.3 provides additional insights based on age and worry levels. It shows that people older than 45 are significantly less likely to perform the measure, with only 10.4% (43 out of 412) intending to do so, compared to 32.4% (183 out of 564) of those younger than 45. The data also indicates that while only a small proportion of the sample (11.6%, 113 out of 976) are worried about the impacts of flooding, these worried individuals show a much higher intention rate of 57.5% (65 out of 113) to implement the measure. In contrast, only 18.7% (161 out of 863) of those who are not worried intend to take action. These figures highlight how factors such as perceived ability, cost perception, age, and worry levels significantly influence the intention to implement this high-effort flood prevention measure.

Figures 6.4 and 6.5 illustrate patterns in intentions to implement a low-effort measure (buying a spare power generator) among people living in the Netherlands. Compared to the high-effort measure, a larger proportion of people intend to implement this measure, with 31.3% (305 out of 976) intending to do so, likely due to its lower effort requirement.

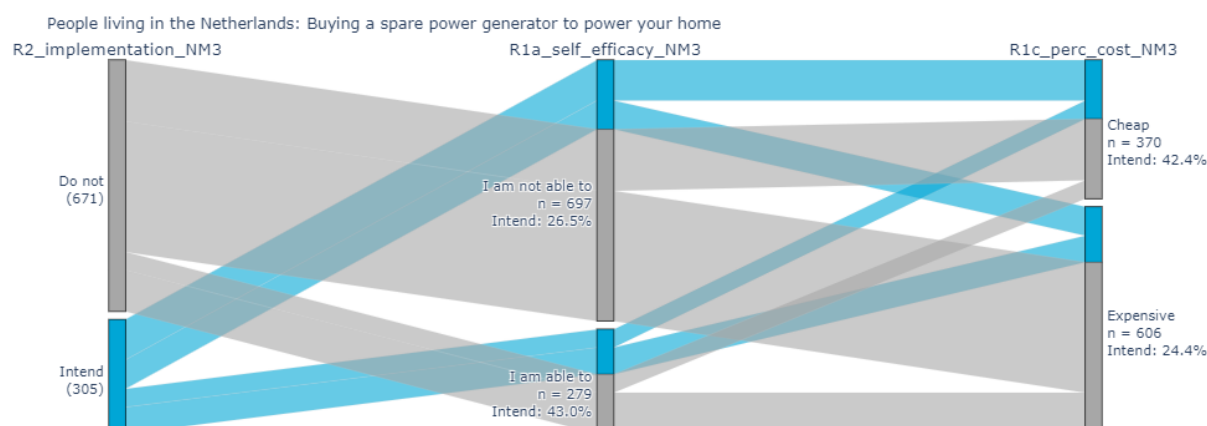


Figure. 6.4. Parallel set diagram for SIRE SCALAR Netherlands analysis - relationships between individuals' intentions to buy a spare power generator (low-effort), self-efficacy and perceived costs.

In Figure 6.4, while 71.4% (697 out of 976) of respondents feel they are not able to implement the measure, 26.5% (185 out of 697) of this group still intend to do so. Among those who feel able to implement it, the intention rate is significantly higher at 43.0% (120 out of 279). Regarding perceived cost, 42.4% (157 out of 370) of those who view the measure as cheap intend to implement it, compared to 24.4% (148 out of 606) of those who perceive it as expensive.

Figure 6.5 shows age and worry levels' influence on intentions. Younger respondents (under 45) show a higher intention rate of 38.3% (216 out of 564) compared to 21.6% (89 out of 412) for those over 45. Worry about flooding impacts also plays a significant role, with 51.1% (58 out of 113) of worried individuals intending to buy a generator, versus only 27.3% (247 out of 863) of those who are not worried.

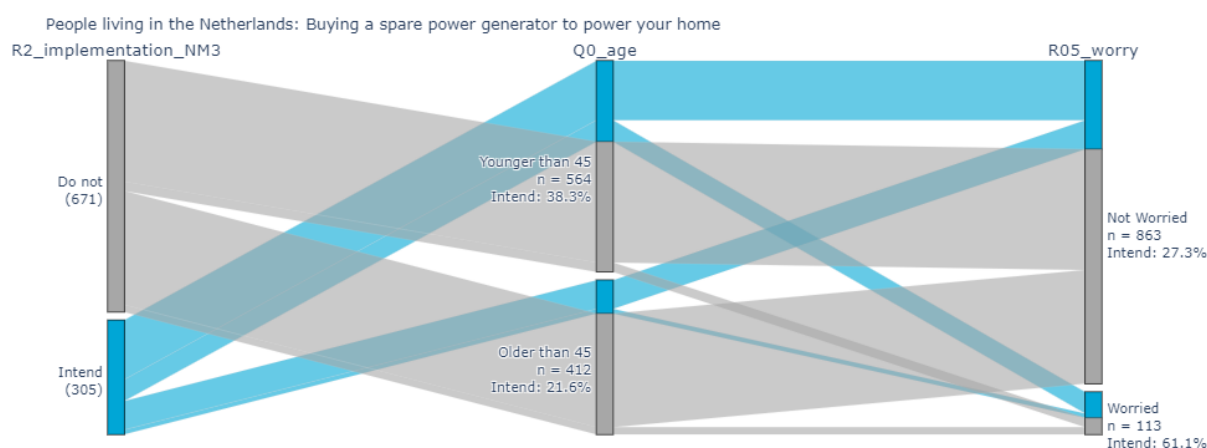


Figure. 6.5. Parallel set diagram for SIRE SCALAR Netherlands analysis - relationships between individuals' intentions to buy a spare power generator (low-effort), age and worry.

While the relationships between conditions and intentions are similar to the high-effort measure, the influence appears weaker. For instance, the difference in intention rates between those who can and cannot implement the measure is smaller (16.5 percentage points) compared to the high-effort scenario (25.8 percentage points). Similarly, the difference based on perceived cost (18.0 percentage points) is less pronounced than in the high-effort case (34.7 percentage points).

This pattern doesn't entirely align with expected outcomes for low-effort measures, where efficacy and perceived cost were anticipated to have a reverse effect. This discrepancy might be attributed to differences in analysis approach and the encoding of results, warranting further investigation into the methodology and data interpretation.

The SIRE method extracts informal rules from the survey data, presenting them in the form of ADICO statements (Attribute, Deontic, aim, Condition, Or else). These statements represent the most significant patterns found in the data, with a focus on the conditions that influence people's intentions regarding flood protection measures. Table 6.1 presents the three lowest entropy extracted informal rules for high-effort measures using the SIRE method on the SCALAR survey data Wave 1 in the Netherlands. A low entropy value implies that the aim response is more homogeneous after being filtered according to the conditions (James, G. et al., 2023). High-effort measures generally have lower entropy which means they have stronger relationships with the conditions overall. All extracted informal rules have aversive aim outcomes, this is because people responded with "do not intend" much more frequently than "do intend" in the survey data. Self-efficacy and cost are common conditions in the informal rules which were expected.

Table. 6.1. The 3 lowest entropy extracted informal rules for high-effort measures using the SIRE method on the SCALAR survey data Wave 1 in the Netherlands

Attribute	Aim	Condition 1	Condition 2	Entropy
People living in the Netherlands	Do not intend to raise the level of the ground floor above the most likely flood level	If they believe that implementing or paying someone to implement this structural measure would be expensive	And if they are not able to undertake the structural measure either themselves or by paying a professional to do so	0.38
People living in the Netherlands	Do not intend to reconstruct or reinforce the walls and/or the ground floor with water-resistant materials	If they are not able to undertake the structural measure either themselves or by paying a professional to do so	And if they believe that implementing or paying someone to implement this structural measure would be expensive	0.40
People living in the Netherlands	Do not intend to strengthen the housing foundations to withstand water pressures	if they are older than 45	if they are not able to undertake the structural measure either themselves or by paying a professional to do so	0.40

These rules provide insights into the conditions most strongly associated with the intention not to implement high-effort flood protection measures. The low entropy values (0.38-0.40) indicate that these rules have a strong relationship with the behaviour. All three rules focus on conditions that lead to not intending to implement measures, reflecting the prevalence of this response in the survey data. The first rule, with the lowest entropy (0.38), suggests that people are unlikely to raise their ground floor level if they perceive it as expensive and are unable to undertake the measure themselves or pay someone else to do it. This aligns with

the visual representation in Figure 6.6 and highlights the importance of both perceived cost and self-efficacy in decision-making.

Figure 6.6 illustrates a Sankey diagram for a high-effort measure: raising the level of the ground floor above the most likely flood level. This diagram visually represents the flow of respondents through different conditions, showing how these factors influence the intention not to implement the measure. The width of each flow corresponds to the number of respondents in that category, allowing for a quick visual comparison of different groups.

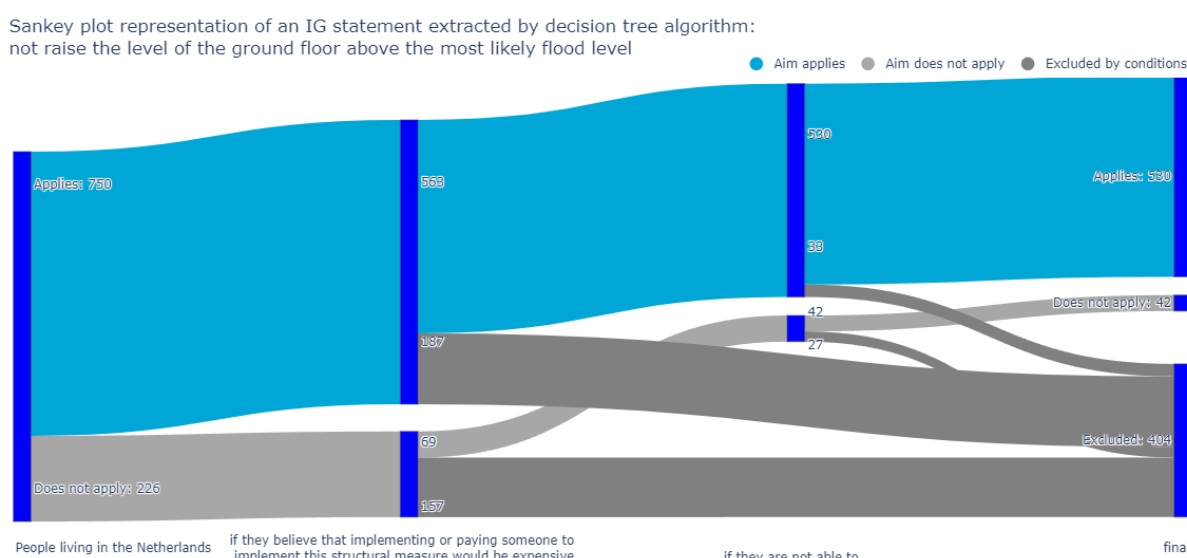


Figure 6.6: Sankey Diagram Illustrating Factors Influencing the Decision to not intend to raise the level of the ground floor above the most likely flood level

In this diagram, out of the total population (976 respondents), 750 do not intend to raise their ground floor level. The two conditions shown - perceiving the measure as expensive and being unable to undertake it - significantly increase the likelihood of not intending to implement the measure. For instance, among those who believe the measure is expensive, a large majority do not intend to implement it. Table 6.2 shows the three lowest entropy extracted informal rules for low-effort measures.

Table. 6.2. The 3 lowest entropy extracted informal rules for low-effort measures using the SIRE method on the SCALAR survey data Wave 1 in the Netherlands

Attribute	Aim	Condition 1	Condition 2	Entropy
People living in the Netherlands	Do not intend to buy a spare power generator to power their home	If they believe that implementing or paying someone to implement this structural measure would be expensive	And if they are not able to buy a spare power generator to power their home	0.68
People living in the Netherlands	Do not intend to install a refuge zone in their home or apartment	if they are not able to install a refuge zone or an opening in the roof of their home or apartment	And if they believe that implementing or paying someone to implement this nonstructural measure would be expensive	0.70
People living in the Netherlands	Do not intend to purchase sandbags or other water barriers	if they are not worried about the potential impact of flooding on their home	if they are older than 45	0.70

The entropy values for low-effort measures (0.68-0.70) are higher than those for high-effort measures, indicating that the conditions have a weaker influence on intentions for low-effort measures. This is consistent with the visual patterns observed in Figure 6.7, where the flows between conditions are less pronounced than in Figure 6.6. Figure 6.7 presents a similar Sankey diagram for a low-effort measure: buying a spare power generator. This diagram shows a higher proportion of people intending to implement the measure compared to the high-effort measure, likely due to its lower perceived difficulty.

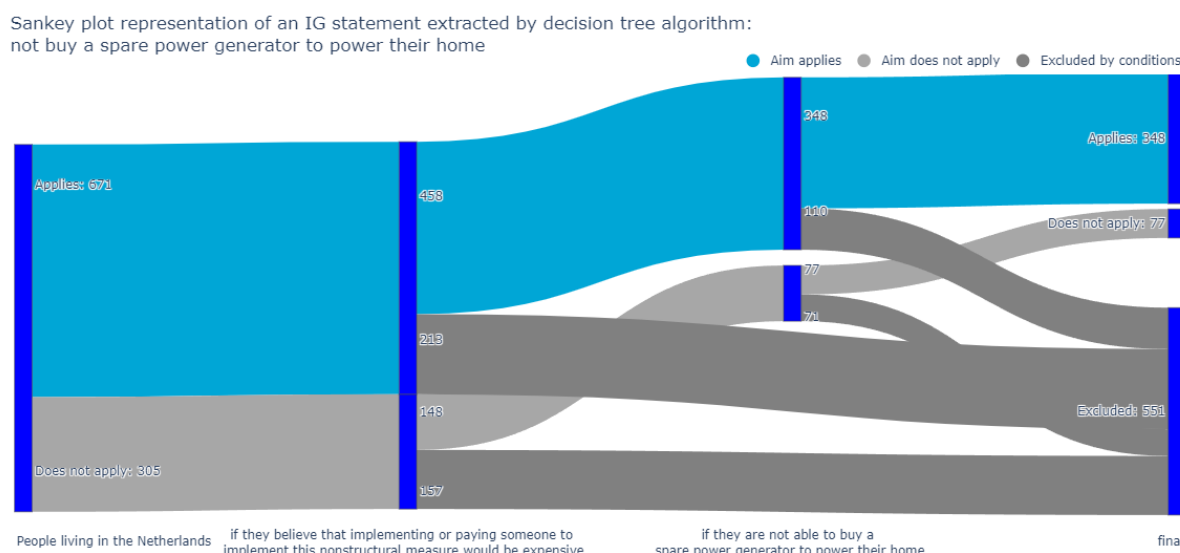


Figure 6.7: Sankey Diagram Illustrating Factors Influencing the Decision to not intend to buy a spare power generator to power their home

Comparing these results to Noll et al. (2022), there are similarities and differences. The SIRE method confirms the importance of perceived cost and self-efficacy, which were also significant in their study. However, the SIRE method also highlights age as a significant factor, with being older than 45 associated with lower intention to implement measures. This difference may be due to the SIRE method's encoding of age into two groups, capturing a trend that wasn't apparent in their regression approach. The ADICO statements extracted from the surveys translate into practical insights for policymakers and flood management professionals. They suggest that interventions to increase flood protection measure adoption should focus on:

1. Reducing perceived costs or providing financial assistance
2. Increasing people's sense of self-efficacy in implementing measures
3. Tailoring approaches for different age groups, with special attention to those over 45
4. Raising awareness about flood risks to increase worry levels, particularly for low-effort measures

It's crucial to recognize, however, that not all policies are translated into formal rules or regulations. From an Institutional Analysis perspective, those policies that do become formalised as regulations are particularly significant. The SIRE method can help bridge the gap between informal rules (as extracted from survey data) and formal rules (regulations) in several ways:

1. Identifying regulatory gaps: By comparing the extracted informal rules with existing regulations, policymakers can identify areas where formal rules are lacking or misaligned with public behaviour and perceptions.
2. Informing regulatory design: The nuanced understanding of factors influencing behaviour (such as age, perceived cost, and self-efficacy) can guide the creation of more targeted and effective regulations.
3. Assessing regulatory impact: The method can be used to evaluate how existing regulations align with or influence informal rules and behaviours, providing insights into the effectiveness of current formal rules.

For example, the finding that perceived cost is a significant barrier could inform the design of regulations that mandate financial assistance programs or tax incentives for flood protection measures. Similarly, the age-related findings might suggest the need for regulations that require tailored outreach and support programs for different age groups.

In conclusion, while the SIRE method largely overlaps with Noll et al. (2022)'s findings, it provides additional nuance and highlights some factors (such as age) that may warrant further investigation. The method's focus on extracting rules for the most common outcomes (in this case, not intending to implement measures) could be seen as a limitation, but it also provides clear direction for where interventions might be most effective. Moreover, by bridging the gap between informal and formal rules, the SIRE method offers a valuable tool for institutional analysts and policymakers to create more effective, responsive, and well-grounded regulations in flood risk management and potentially other domains.

6.2. Validation 2: Policy Analysis and Testing with SCALAR

For the second validation of the SIRE method, a hypothetical situation will be described. A local municipality wants to know if providing a subsidy for households to perform structural changes to their household would be effective (Klijn, F. et al. 2012; KRO-NCRV. 2024; NL Times. 2021). They want to discover what is preventing people from performing these measures and select which measure should receive the subsidy. The two structural changes they are focusing on are:

- **R2_implementation_SM3:** Reconstructing or reinforcing walls and/or the ground floor with water-resistant materials
- **R2_implementation_SM6:** Installing a pump and/or systems to drain flood water.

Then a selection of condition questions are selected to compare the measures and discover factors that might encourage people to implement them.

- **R1a_self_efficacy:** Do you have the ability to undertake the structural measure either yourself or by paying a professional to do so?
- **R01_resilience_6:** My household can rely on the support from my government when I need help (e.g. receiving funding or support in the event of a natural disaster)
- **R08_economic_comfort:** When considering your salary along with your expenses, how would you describe your level of 'economic comfort'?
- **R1c_perc_cost:** When you think in terms of your income and your other expenses, do you believe that implementing or paying someone to implement this structural measure would be cheap or expensive?

The survey responses are then pre-processed to:

Implementation: *Have or will you implement the following measures within the next year?*

1. I have or will implement this measure
2. I will not implement this measure within the next year

Self-efficacy: *Do you have the ability to implement the structural measure?*

1. I am not able
2. I am able

Resilience: *Do you agree that your household can rely on the support from your government when you need help?*

1. I agree
2. I don't agree

Economic Comfort: *How would you describe the level of your salary against your expenses?*

1. *Living comfortably*
2. *Not living comfortably*

Perceived cost: *Do you believe that implementing or paying someone to implement this structural measure would be cheap or expensive?*

1. I believe it would be cheap.
2. I believe it would be expensive.

Following the steps of the SIRE method, parallel set plots were produced using data from the SCALAR Netherlands survey. These plots visualise the relationships between implementation intentions for two structural changes, economic comfort, and perceived government support. The data used for these plots was pre-processed to categorise responses into binary options for each variable, as described in the introduction. Figure 6.8 illustrates the relationships between reconstructing or reinforcing walls and/or ground floor with water-resistant materials (SM3), economic comfort (R08), and perceived government support (R01).

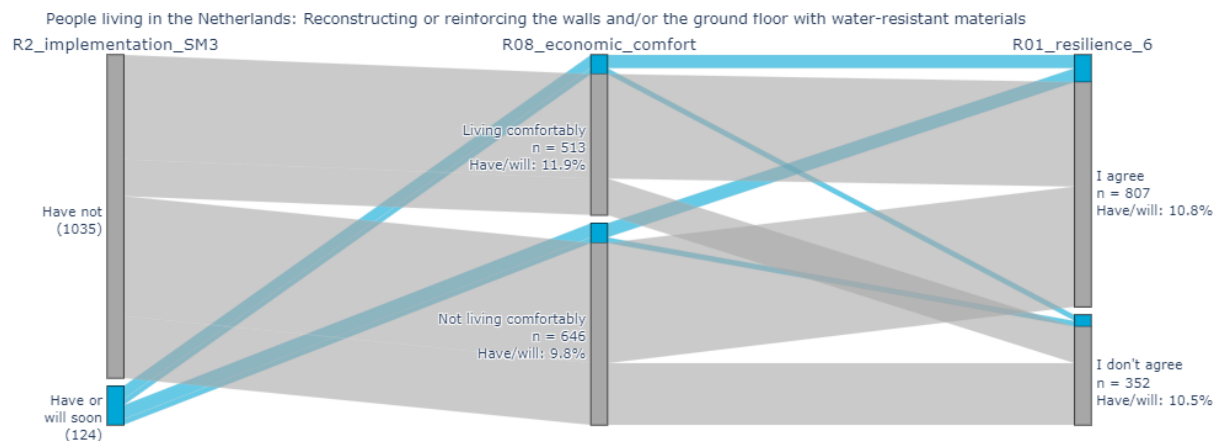


Figure. 6.8. Parallel set diagram for SIRE SCALAR Netherlands analysis - reconstruct or reinforce walls and/or ground floor with water-resistant materials (SM3), economic comfort and resilience through government support.

In Figure 6.8, out of 1,159 total respondents, only 124 (10.8%) have implemented or will reconstruct or reinforce walls and/or ground floor with water-resistant materials within the next year. The economic comfort variable shows a relatively even split, with 513 (44.3%) respondents living comfortably and 646 (55.7%) not living comfortably. Interestingly, those living comfortably show a slightly higher implementation rate (11.9%) compared to those not living comfortably (9.8%). Regarding government support, a majority of 807 (69.6%) agree they can rely on government support, with a similar implementation rate (10.8%) to those who don't agree (10.5% of 352 respondents). Figure 6.9 presents a similar analysis for installing a pump and/or systems to drain flood water (SM6).

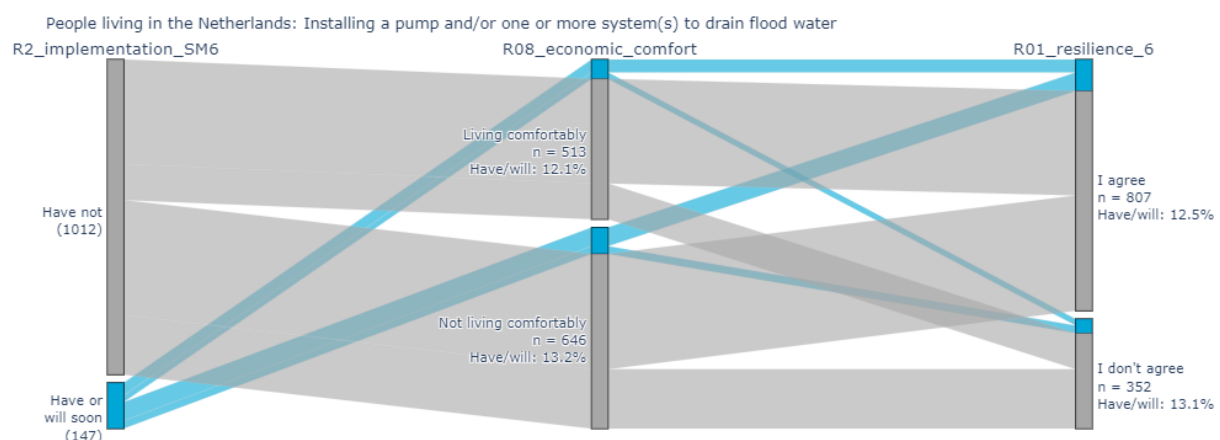


Figure. 6.9. Parallel set diagram for SIRE SCALAR Netherlands analysis - install a pump and/or one or more systems to drain flood water (SM6), economic comfort and resilience through government support.

Figure 6.9 shows that 147 out of 1,159 respondents (12.7%) have implemented or will implement installing a pump and/or one or more systems to drain flood water within the next year, slightly higher than implementing reconstruct or reinforce walls and/or ground floor with water-resistant materials. The economic comfort distribution remains the same as in Figure 6.8. However, the implementation patterns differ slightly. Those not living comfortably show a marginally higher implementation rate (13.2%) compared to those living comfortably (12.1%). For government support, those who agree they can rely on support have a slightly lower implementation rate (12.5%) compared to those who don't agree (13.1%).

A comparison of the two figures reveals that installing pumps or drainage systems has a slightly higher overall implementation rate than reconstructing with water-resistant materials. This might suggest that people find installing a pump or systems to drain flood water more feasible or necessary. In both cases, the majority of respondents (69.6%) agree that they can rely on government support. However, this perceived support doesn't seem to strongly influence implementation intentions for either measure. The relationship between economic comfort and implementation intentions is inconsistent between the two measures, with comfortable respondents slightly more likely to implement reconstructing or reinforcing walls and/or ground floor with water-resistant materials, but slightly less likely to implement installing a pump or systems to drain flood water.

These observations suggest that while economic factors and perceived government support play a role in implementation decisions, their influence is not straightforward and may vary depending on the specific measure. The low overall implementation rates (10.7% for SM3 and 12.7% for installing a pump or systems to drain flood water) indicate that other factors not captured in these variables may be more significant in determining whether people implement these flood protection measures.

Figures 6.10 and 6.11 present parallel set diagrams that examine the relationships between self-efficacy, perceived costs, and the implementation of two structural measures for flood protection in the Netherlands. These diagrams are based on data from the SCALAR Netherlands survey, processed according to the SIRE method. Figure 6.10 focuses on reconstructing or reinforcing walls and/or ground floor with water-resistant materials (SM3).

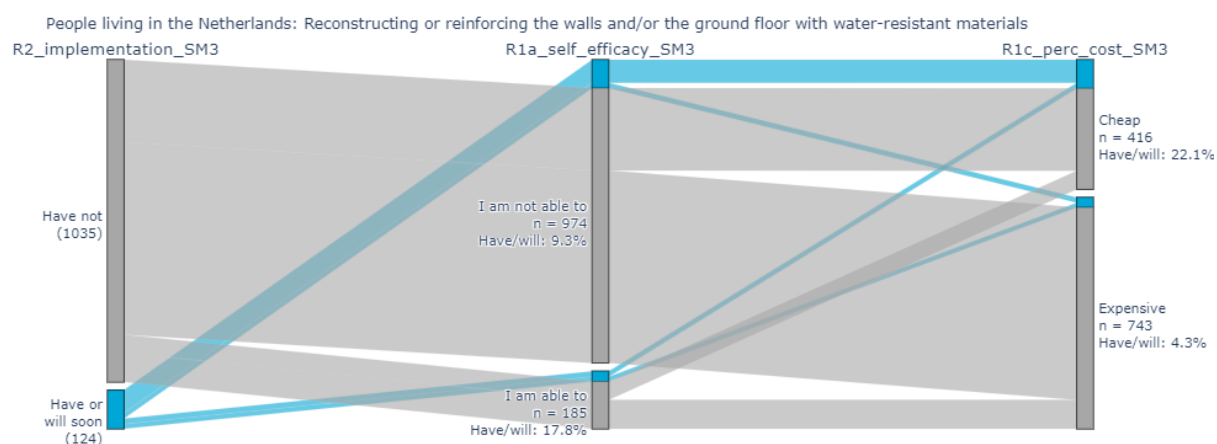


Figure 6.10. Parallel set diagram for SIRE SCALAR Netherlands analysis - reconstruct or reinforce walls and/or ground floor with water-resistant materials (SM3), self-efficacy and perceived costs.

Figure 6.10 reveals that out of 1,159 total respondents, 124 (10.7%) have implemented or

will implement this measure within the next year. Regarding self-efficacy, only 185 respondents (16.0%) consider themselves able to implement SM3, while 974 (84.0%) do not. Interestingly, among those who have implemented or will implement SM3, 26.6% (33 out of 124) consider themselves able to do so, while 63.4% (91 out of 124) do not. In terms of perceived costs, 416 respondents (35.9%) consider SM3 cheap, while 743 (64.1%) perceive it as expensive. The implementation rate is significantly higher among those who perceive reconstructing walls and/or ground floor with water-resistant materials as cheap (22.1%, or 92 out of 416) compared to those who perceive it as expensive (4.3%, or 32 out of 743). Figure 6.11 presents a similar analysis for installing a pump and/or systems to drain flood water.

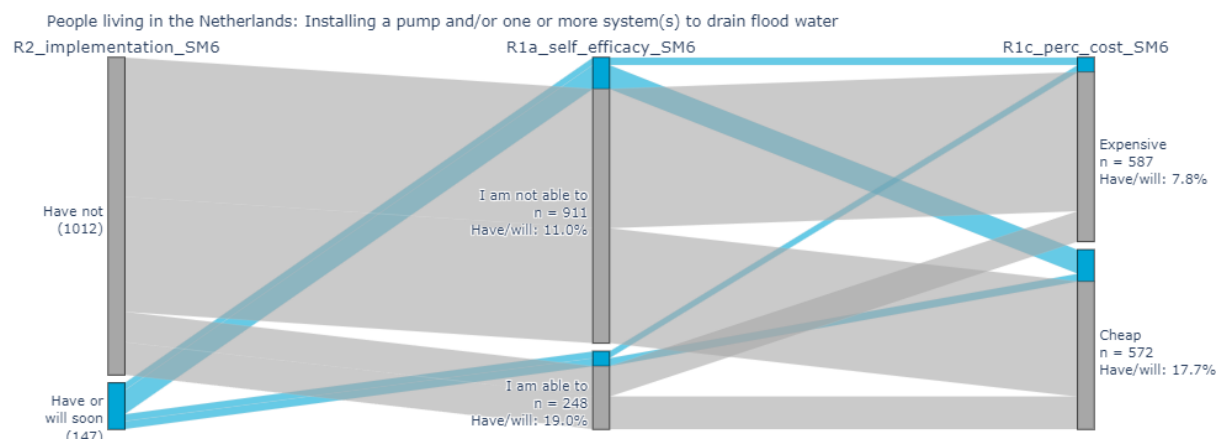


Figure. 6.11. Parallel set diagram for SIRE SCALAR Netherlands analysis - install a pump and/or one or more systems to drain flood water (SM6), self-efficacy and perceived costs.

Figure 6.11 shows that out of 1,159 respondents, 147 (12.7%) have implemented or will implement this measure within the next year. A higher number of respondents (248, or 21.4%) consider themselves able to implement installing a pump or systems to drain flood water compared to reconstructing walls and/or floor with water-resistant materials, while 911 (78.6%) do not. Among those who have implemented or will implement installing a pump or systems to drain flood water, 32.0% (47 out of 147) consider themselves able to do so, while 68.0% (100 out of 147) do not. Regarding perceived costs, 572 respondents (49.4%) consider installing a pump or systems to drain flood water cheap, while 587 (50.6%) perceive it as expensive. The implementation rate is higher among those who perceive installing a pump or systems to drain flood water as cheap (17.7%, or 101 out of 572) compared to those who perceive it as expensive (7.8%, or 46 out of 587).

Comparing the two figures reveals that more respondents consider themselves able to install a pump or systems to drain flood water (21.4%) than reconstruct walls and/or ground floor with water-resistant materials (16.0%). Similarly, a larger proportion of respondents perceive installing a pump or systems to drain flood water as cheap (49.4%) compared to reconstructing walls and/or ground floor with water-resistant materials (35.9%). These differences may contribute to the slightly higher overall implementation rate for installing a pump or systems to drain flood water (12.7%) compared to SM3 (10.7%).

In both cases, perceived cost appears to be a strong factor influencing implementation intentions. Respondents who perceive the measures as cheap are significantly more likely to implement them. However, the relationship between self-efficacy and implementation is less

straightforward. A notable proportion of respondents who have implemented or will implement these measures do not consider themselves able to do so. This unexpected result could be due to respondents working towards making the measure possible within the next year, or having already implemented the measure and no longer considering it a future possibility. These observations suggest that while perceived cost strongly influences implementation decisions for both measures, the role of self-efficacy is more complex and may require further investigation to fully understand its impact on flood protection measure implementation.

The SIRE method extracts informal rules from the survey data, presenting them in Table 6.3, which shows the five least entropy informal rule statements for the selected measures. A low entropy value indicates that the aim response is more homogeneous after being filtered according to the conditions (James, G. et al. 2023). Table 6.3 presents the extracted informal rules, their conditions, and their corresponding entropy values. The table reveals important insights into factors influencing the implementation of flood protection measures in the Netherlands.

Table. 6.3. The 5 lowest entropy extracted informal rules using the SIRE method for household structural changes: reconstruct or reinforce walls and/or ground floor with water-resistant materials, install a pump and/or one or more systems to drain flood water

Attribute	Aim	Condition 1	Condition 2	Entropy
People living in the Netherlands	Will not reconstruct or reinforce walls and/or ground floor with water-resistant materials within the next year	if they believe that implementing or paying someone to implement this structural measure would be expensive	And if they are not able to undertake the structural measure either themselves or by paying a professional to do so	0.18
People living in the Netherlands	Will not reconstruct or reinforce walls and/or ground floor with water-resistant materials within the next year	if they believe that implementing or paying someone to implement this structural measure would be expensive	None	0.26
People living in the Netherlands	Will not install a pump and/or one or more systems to drain flood water within the next year	if they believe that implementing or paying someone to implement this structural measure would be expensive	And if they do not have the ability to undertake the structural measure either themselves or by paying a professional to do so	0.29
People living in the Netherlands	Will not install a pump and/or one or more systems to drain flood water within the next year	if they believe that implementing or paying someone to implement this structural measure would be expensive	None	0.40
People living in the Netherlands	not reconstruct or reinforce walls and/or ground floor with water-resistant materials within the next year	if they are not able to undertake the structural measure either themselves or by paying a professional to do so	if they feel they cannot rely on government support	0.44

The extracted informal rule statements in Table 6.3 indicate that reconstructing or reinforcing walls and/or ground floor with water-resistant materials (SM3) is less likely to be implemented within the next year compared to installing a pump and/or systems to drain flood water (SM6). This is evidenced by the lower entropy values for SM3 (0.18 and 0.26) compared to SM6 (0.29 and 0.40), suggesting that the conditions for not implementing SM3 are more consistently met.

The perception of high costs emerges as a strong deterrent for both measures, appearing as a condition in four out of the five rules. This suggests that policies aimed at increasing adoption should address financial barriers through subsidies or financial incentives.

Additionally, the ability to personally undertake the measure or hire a professional is a significant factor, especially when combined with cost concerns. This is evident in the lowest entropy rule (0.18) for SM3, which combines both cost and ability conditions. Providing technical support or facilitating access to professionals could help overcome this barrier.

The extracted informal rules can be visually represented in Sankey diagrams. Figures 6.12, 6.13, and 6.14 provide Sankey diagrams illustrating the factors influencing the decision not to implement flood water resistance or drainage systems in the Netherlands. These diagrams visually represent the flow of respondents through different conditions, showing how these factors influence the intention not to implement the measures. In Sankey diagrams, the width of the flows represents the number of respondents, and the colours indicate different categories or outcomes.

Figure 6.12 focuses on factors influencing the decision not to reconstruct or reinforce walls and/or ground floor with water-resistant materials within the next year.

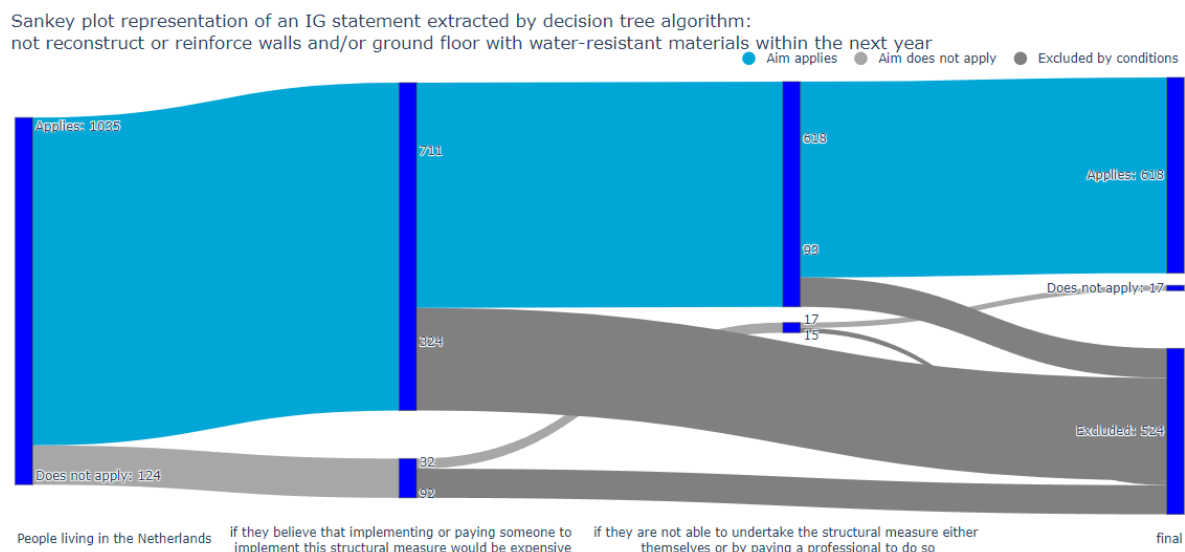


Figure 6.12: Sankey Diagram Illustrating the First Row of Table 6.3, Factors Influencing the Decision to Not Reconstruct Walls/Floor with Water-Resistant Materials in the Netherlands.

In Figure 6.12, we can see that out of the total population (1159 respondents), a significant majority (1035) indicate they will not implement the measure. The first condition, believing that implementing the measure would be expensive, excludes 711 respondents from potentially implementing the measure. The second condition, not being able to undertake the measure, further reduces the number of potential implementers to just 47. The diagram clearly shows how these two conditions dramatically increase the proportion of those not intending to implement the measure, with only 47 respondents still considering implementation after both conditions. Figure 6.13 presents a similar pattern but with only one condition.

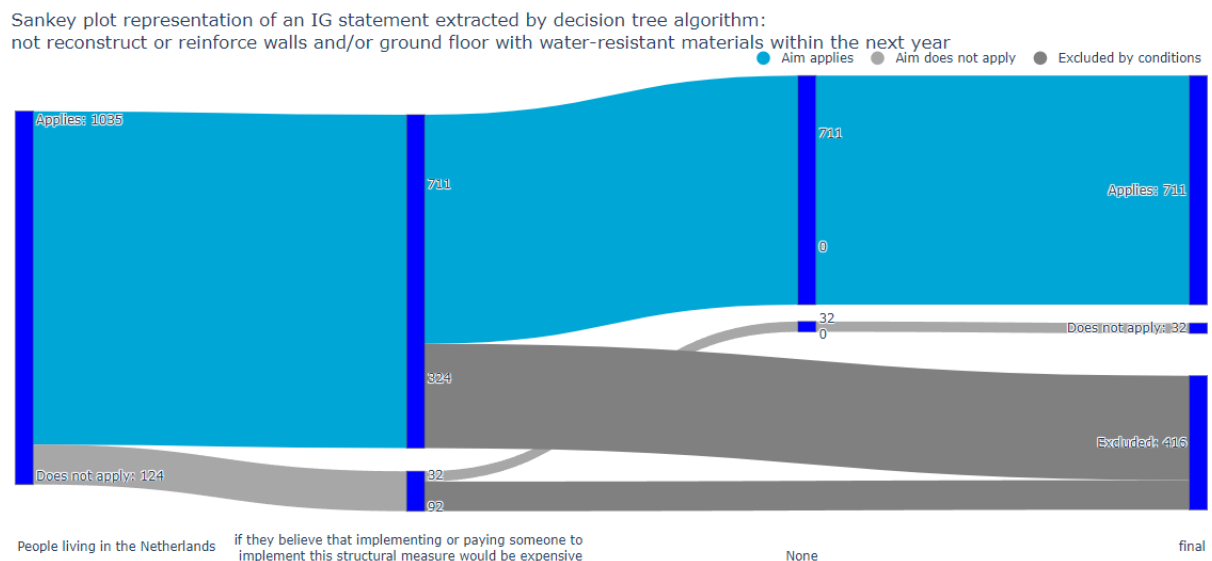


Figure 6.13: Sankey Diagram Illustrating the Second Row of Table 6.3, Factors Influencing the Decision to Not Reconstruct Walls/Floor with Water-Resistant Materials in the Netherlands.

Figure 6.13 illustrates the impact of perceiving the measure as expensive on the decision not to reconstruct or reinforce walls/floor with water-resistant materials. Out of 1159 respondents, 1035 initially do not intend to implement the measure. After applying the condition of perceived expense, 711 respondents are excluded, leaving only 324 still considering implementation. This reinforces the strong influence of perceived cost on implementation decisions. Figure 6.14 shifts focus to factors influencing the decision not to install flood water drainage systems.

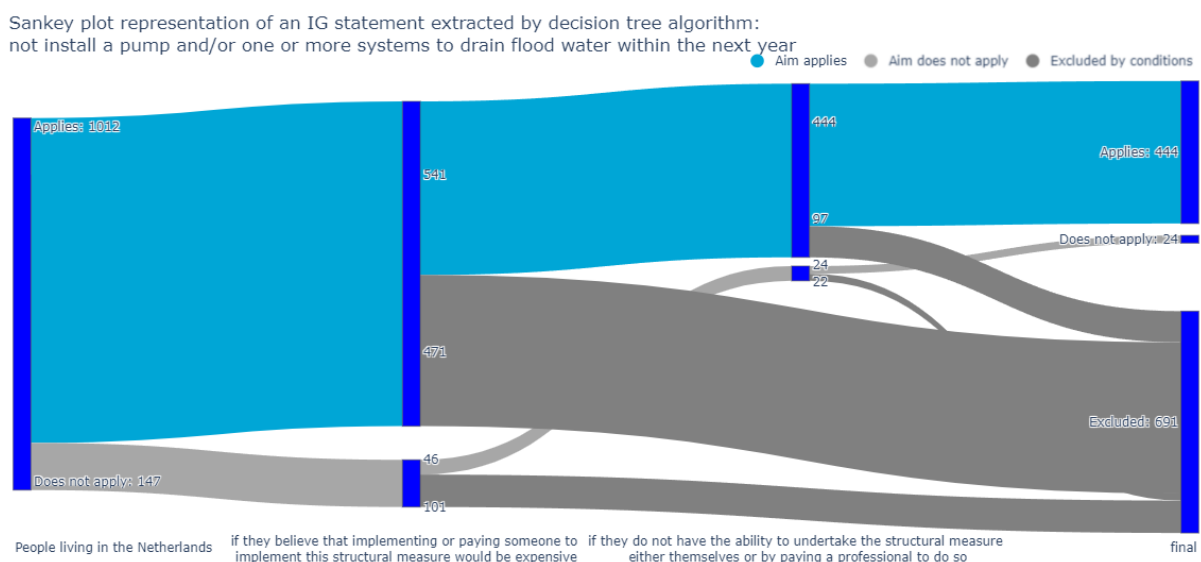


Figure 6.14: Sankey Diagram Illustrating the Third Row of Table 6.3, Factors Influencing the Decision to Not Install Flood Water Drainage Systems in the Netherlands.

Figure 6.14 shows a slightly different pattern compared to the previous two. Out of 1159 respondents, 1012 initially do not intend to install drainage systems. The first condition (perceiving the measure as expensive) excludes 588 respondents, while the second condition (lack of ability to undertake the measure) further reduces potential implementers to 449. The slightly higher number of respondents still considering implementation after both

conditions (449 compared to 47 in Figure 6.12) suggests that installing drainage systems might be perceived as more feasible or necessary than reconstructing walls/floors.

These Sankey diagrams effectively reinforce the findings from Table 6.3, visually demonstrating the strong influence of perceived cost and self-efficacy on implementation decisions. They provide a clear, intuitive representation of how different factors progressively influence decision-making, offering valuable insights for policy development and targeted interventions.

This validation procedure demonstrates a potential use case for the SIRE method. While the findings are relatively limited due to the selected questions and encoding choices, and the aims already having relatively homogenous responses, the results are informative and produce logical policy advice. The analysis suggests that effective strategies for increasing adoption of flood protection measures should include offering financial assistance and technical support. Initially focusing on measures with higher perceived efficacy and lower costs could achieve quicker adoption and demonstrate success, potentially paving the way for more extensive measures such as installing a pump or systems to drain flood water. The method's ability to extract meaningful insights from survey data and present them in both tabular and visual formats highlights its potential as a valuable tool for institutional analysis.

The SIRE method has shown several strengths in this validation. The use of parallel set diagrams and Sankey diagrams offers intuitive visualisations of complex relationships between variables, enhancing interpretability. The method's ability to calculate entropy values for extracted rules provides a measure of their significance and reliability. The extracted informal rules directly inform potential policy interventions, bridging the gap between data analysis and practical application.

However, the application also revealed areas for improvement. The method's results can be influenced by how variables are categorised, as seen in the age-related findings. This highlights the need for careful consideration in data preprocessing. The analysis focused on a few key variables. Expanding the range of conditions examined could provide a more comprehensive understanding of decision-making factors. The method may also be less effective when dealing with highly homogenous responses, potentially missing nuanced variations in behaviour.

In conclusion, while the SIRE method shows promise as a tool for institutional analysis in the context of flood protection measures, its effectiveness can be further enhanced. By addressing its current limitations and building on its strengths, the SIRE method has the potential to become a robust and widely applicable approach for extracting informal rules from survey data, thereby informing more effective and targeted policy interventions in complex social-ecological systems.

6.3. Validation 3: Policy Analysis and Testing with ESS

For the third validation approach, the SIRE method is implemented as a strategic analytics tool for political voting behaviour analysis. This application aims to demonstrate the method's versatility and its potential use in policy development and political strategy. The goal is to identify conditions that increase the likelihood of people voting for left-leaning parties in the

Netherlands and to derive corresponding policy or marketing strategies for these parties. The European Social Survey (ESS) provides a rich dataset for this analysis, including over 500 questions on various topics such as political opinions and attitudes towards institutional bodies. Unlike the SCALAR dataset used in previous validations, which focused on specific behaviours and their causes, the ESS emphasises broader trends and patterns in social and political attitudes.

A key question in the ESS asks respondents from the Netherlands which party they voted for in the last national election (variable: prtvtthnl), note that this survey was performed in 2020 and therefore these results refer to two election cycles ago. This question will serve as the Aim to analyse. For this analysis, five left-leaning parties are selected as a grouped Aim outcome (Otjes, S. 2018; Kuitert, G. 2021):

1. Labour Party
2. Socialist Party
3. Democrats '66
4. Green Left
5. Volt

Responses to this question will be encoded into two categories: "voted for selected left-leaning parties" or "voted for other party." Responses such as "Not applicable," "Refusal," "Don't know," or "No answer" will be excluded from the analysis to ensure focus on actual voting behaviour. This application of the SIRE method to voting behaviour presents unique challenges and opportunities. It requires a different approach compared to the previous validations. The analysis will leverage the SIRE method to extract informal rules related to voting behaviour, potentially uncovering insights that could inform political strategies and policy development for left-leaning parties in the Netherlands.

First, parallel set diagrams are generated to observe how the responses to the voting behaviour question distribute amongst the responses to other questions. Figures 6.15 and 6.16 below show examples of parallel set diagrams that were examined to better understand the survey data responses. As can be seen in Figures 6.15 and 6.16, many questions in ESS such as these, have far too many response options, showing the need for encoding or grouping responses to produce better visualisations and discover relationships.

Language most often spoken at home: first mentioned

Voted for?

Region	n	Voted Left (%)	Language(s) Spoken at Home	Voted for a left-leaning party (n=497)	Voted for another party (n=618)
Zuid-Holland	195	37.4%	Dutch	Yes	No
Overijssel	96	36.5%	Dutch	Yes	No
Groningen	42	57.1%	Dutch	Yes	No
Noord-Holland	161	55.9%	Dutch	Yes	No
Gelderland	144	43.8%	Dutch	Yes	No
Limburg (NL)	61	37.7%	Dutch	Yes	No
Utrecht	129	52.7%	Dutch	Yes	No
Noord-Brabant	159	47.8%	Dutch	Yes	No
Drenthe	30	23.3%	Dutch	Yes	No
Friesland	26	23.1%	Dutch	Yes	No
Friesland (NL)	49	-	Dutch	Yes	No
Zeeland	23	34.8%	Dutch	Yes	No
Total	1067	44.7%	Dutch	497	618

Language most often spoken at home: second mentioned

- No answer
- Western Frisian
- Don't know
- English
- Turkish
- Arabic
- Berber (Other)
- Kurdish
- Spanish/Castilian
- Chinese
- Other

Voting Behavior Summary:

- Voted for a left-leaning party: 497
- Voted for another party: 618

Figure 6.15 presents a parallel set diagram illustrating the distribution of voting behaviour across regions and languages spoken at home in the Netherlands. The diagram shows that voting patterns for left-leaning parties vary considerably across regions, with some areas like Groningen showing a higher proportion of left-leaning voters (57.1%) compared to others like Limburg (35.7%). The language spoken at home appears to have some correlation with voting behaviour, though the relationship is not straightforward. This visualisation highlights the complexity of voting patterns and suggests that regional and linguistic factors may play a role in political preferences.

People: Party voted for in last national election, Netherlands

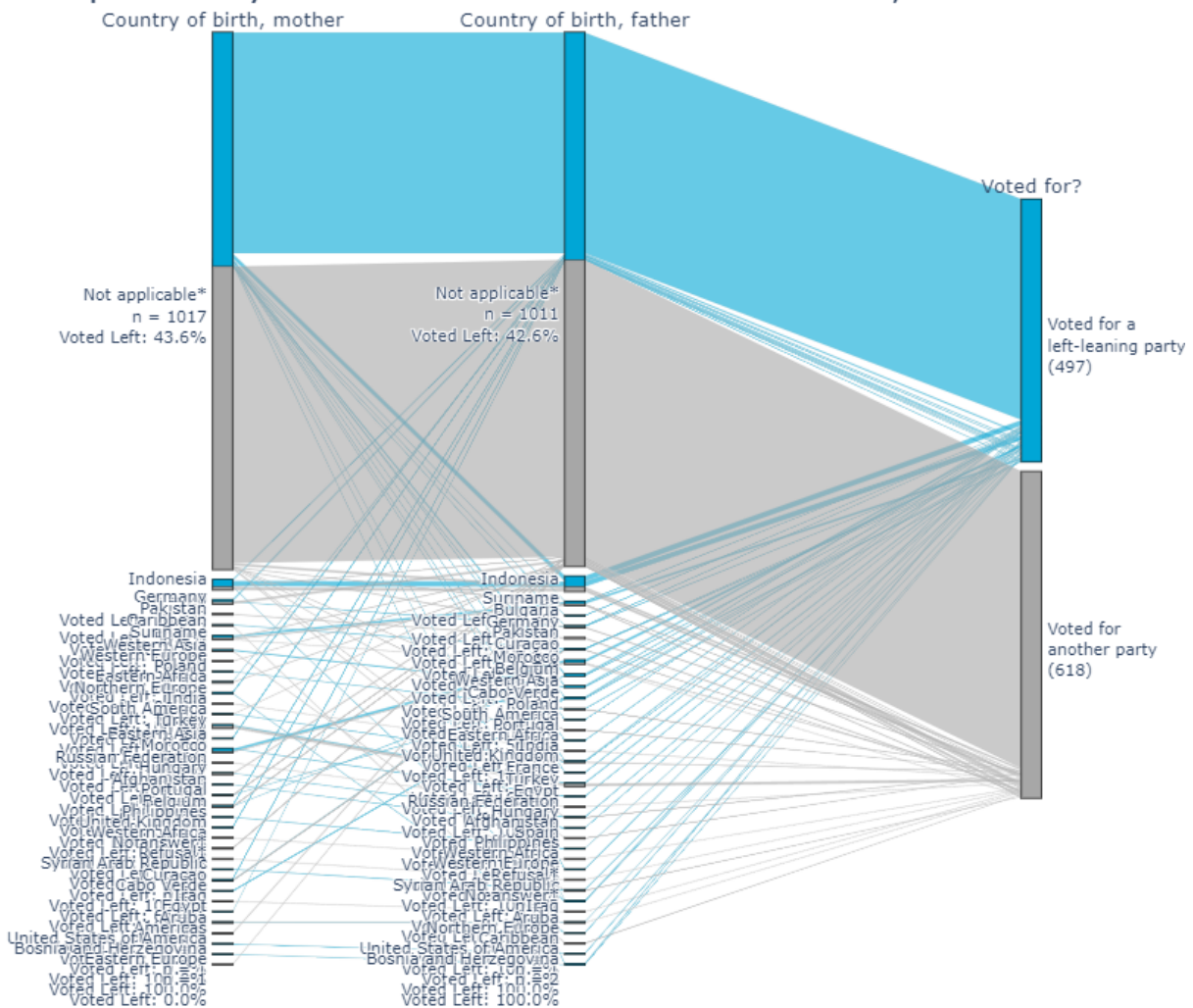


Figure 6.16: Parallel Set Diagram showing example response distribution of voting behaviour amongst other questions in the ESS Data.

Figure 6.16 displays the relationship between parents' country of birth and voting behaviour. The diagram reveals that a large majority of respondents' parents were born in the Netherlands. Interestingly, among those whose parents were born outside the Netherlands, there seems to be a slightly higher tendency to vote for left-leaning parties. This suggests that immigrant background might have some influence on political preferences, although the effect appears to be modest.

Instead of encoding all questions, the decision tree algorithm is first applied to analyse the ESS responses as described in the SIRE method. The analysis identifies a set of conditions that increase voting outcome homogeneity. The conditions splitting the responses in the decision trees will be compiled to create a selection of ESS questions that are further analysed and encoded. Figures 6.17 and 6.18 present two of the three decision trees produced in the first round of analysis to select condition questions related to voting behaviour. The tree is based on an entropy threshold of 0.85 and a sample size threshold of 0.3, any nodes that meet these requirements will be used to form informal rules. It's important to note that responses such as "don't know" and "no response" have been

included, which may affect the relationship between the aim question and ordered categorical questions as these responses don't fit within the range of response options.

In Figures 6.17 and 6.18, the first round of decision trees are shown. The root node shows a relatively balanced split between those who voted for the selected left-leaning parties (528 samples) and those who voted for other parties (661 samples), the entropy is close to 1 with 0.991.

prtvthnl: Party voted for in last national election, Netherlands

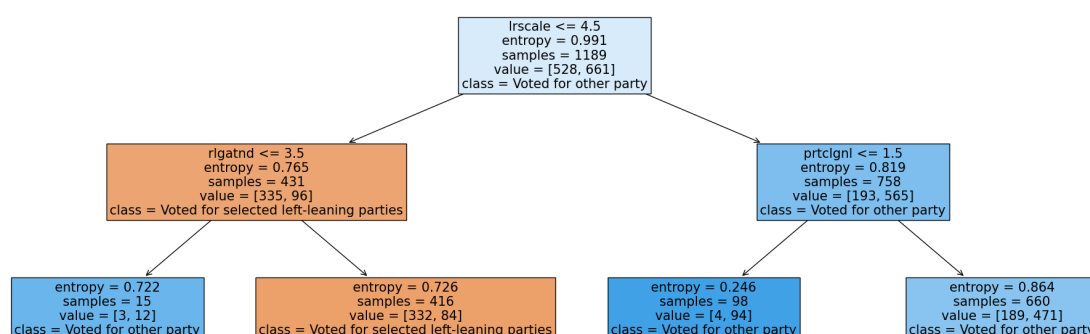


Figure 6.17: Example decision tree produced in the first round to select a set of condition questions with relationships to the voting behaviour aim where all questions from the ESS data are included.

In Figure 6.17, the first split is based on the "lrscale" variable, which relates to whether the responder considers themselves more left or right leaning. The node on the left of the root node contains responders that consider themselves more left leaning, with a sample size of 431 and a lower entropy of 0.765 with 335 responders at the node having voted for the selected parties. The node on the right of the root node contains responders that consider themselves more right leaning, with a sample size of 758 and a slightly higher entropy of 0.819 with 565 responders at the node having voted for the parties that were not selected. This could be due to the selection of parties (e.g. some of the selected parties could have right leaning policies or there were parties that were not selected and could be considered left), or due to people choosing to vote for a party that leans more left than their usual ideology. The responses on the left node are subsequently split based on the "rlgatnd" variable, which relates to how often they attend religious services apart from special occasions. This suggests that a relation exists between infrequently attending religious services and voting for the selected left leaning parties.

prtvthnl: Party voted for in last national election, Netherlands

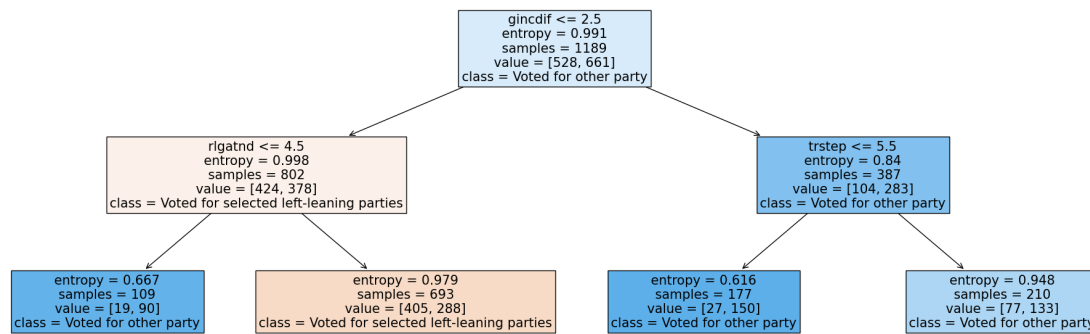


Figure 6.18: Example decision tree produced in the first round to select a set of condition questions with relationships to the voting behaviour aim where all questions from the ESS data are included.

In Figure 6.18, the first split is based on the "gincdif" variable, which relates to opinions on government reducing income differences. The node on the left contains the responders that believe the government should reduce income differences with a sample size of 802 and an increased entropy of 0.998. This implies that a majority of voters, regardless of their voting choice, believe that the government should reduce income inequality. The node on the right has 387 samples, an entropy of 0.84, and represents people that believe that the government does not need to deal with income differences. This node is the only node that is selected by the SIRE method as it is the only node that meets the entropy and sample size thresholds. The subsequent nodes below it are not selected by the SIRE method to form informal rules because the sample size does not meet the set threshold of 0.3.

This experimental approach using the decision tree to select questions for encoding and cleaning, rather than cleaning all questions upfront, allows for a focused analysis on the more relevant variables. However, it's important to recognize that this method may miss some nuanced relationships and could be influenced by the initial data structure. Future refinements could involve iterative cleaning and tree-building processes to capture a broader range of relevant variables. The selected questions in this instance are:

- **gincdif**: Government should reduce differences in income levels
- **rlgatnd**: How often attend religious services apart from special occasions
- **hmsacld**: Gay and lesbian couples right to adopt children
- **keydec**: Key decisions are made by national governments rather than the European Union
- **Irscale**: Placement on left right scale

These selected questions are then encoded as follows:

Government should reduce differences in income levels (gincdif):

1. Agree
2. Do not agree

How often attend religious services apart from special occasions (rlgatnd):

1. Attends religious services
2. Rarely attends religious services

Placement on left right scale (lrscale):

1. Consider themselves more left
2. Consider themselves more right

Key decisions are made by national governments rather than the European Union (keydec):

1. Not important for democracy
2. Important for democracy

Gay and lesbian couples right to adopt children (hmsacld):

1. Agree
2. Do not agree

This encoding process simplifies the response options, making them more suitable for analysis and visualisation. It focuses on the extremes of each scale, which can help highlight clear distinctions in attitudes and their relationship to voting behaviour. After encoding, these questions can then be visualised as parallel set diagrams to ensure they are numerically and contextually relevant to the aim. These visualisations are shown in Figures 6.19 to 6.21.

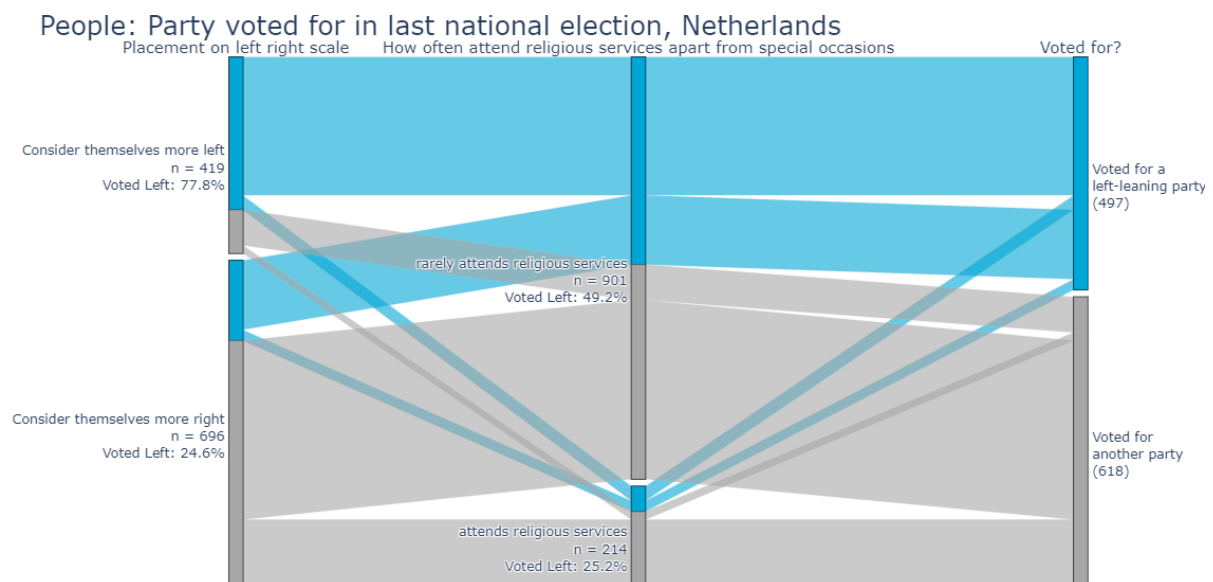


Figure 6.19: Parallel Set Diagram with decision tree selected conditions that have been encoded to show response distribution of voting behaviour in the ESS Data.

Figure 6.19 presents a parallel set diagram illustrating the relationship between left-right political alignment, religious service attendance, and voting behaviour. Of those who consider themselves more left (419 respondents), a significant majority (77.8%) voted for left-leaning parties. Conversely, among those who consider themselves more right (696 respondents), only 24.6% voted for left-leaning parties. Religious service attendance also shows a pattern: those who rarely attend religious services (901 respondents) are more likely to vote for left-leaning parties (47.7%) compared to those who regularly attend (214 respondents), of whom only 25.2% voted for left-leaning parties.

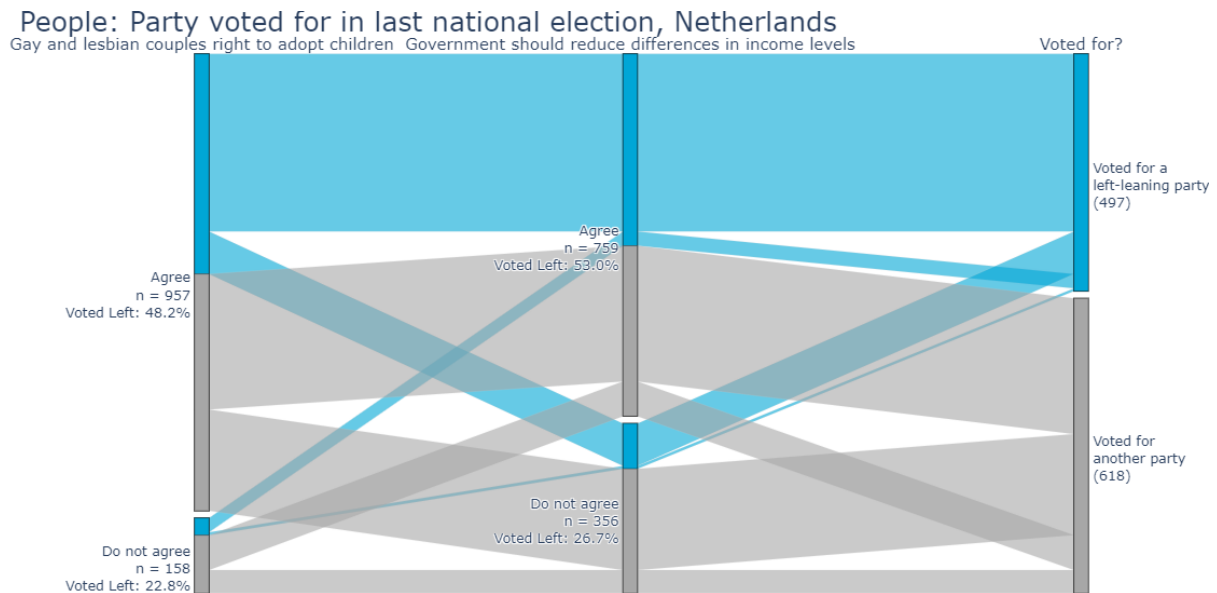


Figure 6.20: Parallel Set Diagram with decision tree selected conditions that have been encoded to show response distribution of voting behaviour in the ESS Data.

Figure 6.20 demonstrates the relationship between attitudes towards gay and lesbian couples' right to adopt children, views on government reducing income differences, and voting behaviour. Those who agree with adoption rights for gay and lesbian couples (957 respondents) are more likely to vote for left-leaning parties (48.2%) compared to those who disagree (158 respondents), of whom only 22.8% voted for left-leaning parties. Similarly, those who agree that the government should reduce income differences (769 respondents) are more likely to vote for left-leaning parties (53.6%) compared to those who disagree (336 respondents), of whom only 26.5% voted for left-leaning parties.

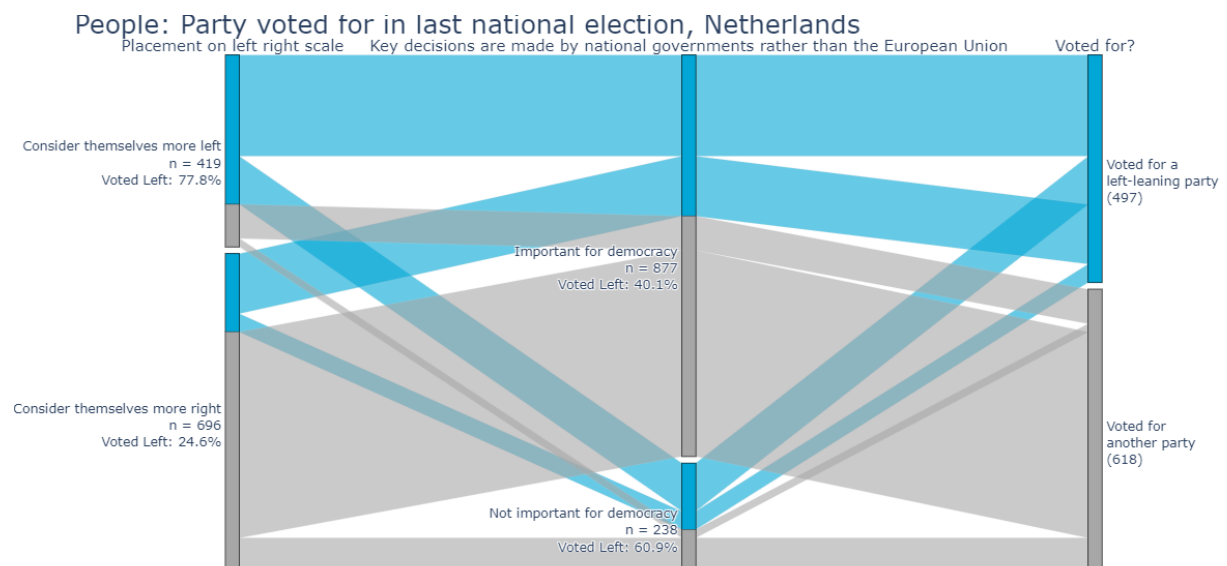


Figure 6.21: Parallel Set Diagram with decision tree selected conditions that have been encoded to show response distribution of voting behaviour in the ESS Data.

Figure 6.21 shows the relationship between left-right political alignment, views on key decision-making by national governments versus the EU, and voting behaviour. As seen in Figure 6.19, those who consider themselves more left are significantly more likely to vote for

left-leaning parties. Interestingly, among those who consider key decisions by national governments as not important for democracy (238 respondents), a higher proportion (60.9%) voted for left-leaning parties, compared to those who consider it important (877 respondents), of whom 40.7% voted for left-leaning parties. It is immediately apparent that these figures are much cleaner and patterns are much clearer than in the previous parallel set diagram figures 6.15 and 6.16. This improvement in clarity demonstrates the value of the encoding process in simplifying complex data for more effective analysis.

After encoding these selected questions and visualising them in parallel set diagrams, the decision tree algorithm is repeated with only these questions to see how the relationship changes and what informal rules emerge. In this instance, the sample size limit is set to 0.33 and entropy to 0.95. This produced decision trees such as Figures 6.22 and 6.23.

prtvthnl: Party voted for in last national election, Netherlands

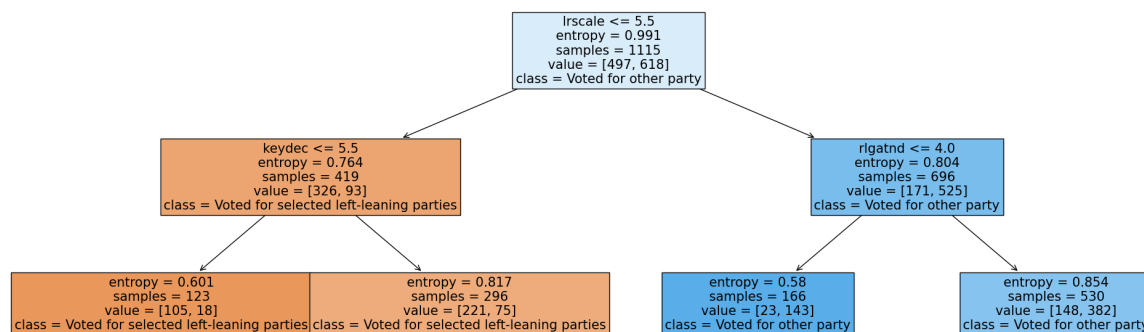


Figure 6.22: Example decision tree produced in the second round after encoding the selected condition questions with relationships to the voting behaviour aim where only the selected questions are included.

Figure 6.22 shows a decision tree with the root node split based on the 'Irscale' variable (left-right scale). For those who consider themselves more left (419 samples), the majority (326) voted for selected left-leaning parties. For those who consider themselves more right (696 samples), the majority (525) voted for other parties. The tree further splits based on the 'rlgatnd' variable (religious service attendance) for those on the right, and 'keydec' (opinion on national governments making key decisions versus European Union) for those on the left.

prtvthnl: Party voted for in last national election, Netherlands

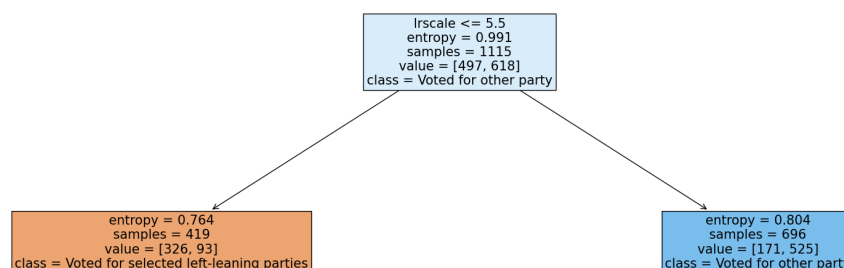


Figure 6.23: Example decision tree produced in the second round after encoding the selected condition questions with relationships to the voting behaviour aim where only the selected questions are included.

Figure 6.23 presents a decision tree with only one split based on the 'Irscale' variable, none of the other questions could be used to split the responses and form a node with a sample size and entropy that meet the thresholds. This tree suggests that the left-right political alignment is the most significant predictor of voting behaviour among the encoded variables. The two informal rules extracted in this instance are shown in Table 6.4.

Table 6.4: Table showing extracted informal rules when entropy threshold is 0.95 and sample size threshold set to 0.33.

Attribute	Aim	Condition 1	Condition 2	Entropy
People	Voted for selected left-leaning parties	if they consider themselves more left	And none	0.76
People	Voted for other party	if they consider themselves more right	And none	0.80

It's important to note that this outcome is not surprising. Only one question (left-right scale) has been used by the decision tree algorithm because the other encoded questions do not split the data such that it meets the entropy and sample size requirement. While this might suggest flaws in the approach to select questions, it also validates that the SIRE method extracts logical outcomes. Other reasons for the weaker-than-expected relationships could include, the oversimplification of complex political attitudes through binary encoding or the selected left-leaning parties may not fully represent the left spectrum.

This shows the importance for researchers using the SIRE method to always conduct preliminary research to understand which questions they want to test. These extracted informal rules are represented in the Sankey plots for Figures 6.24 and 6.25.

Sankey plot representation of an IG statement extracted by decision tree algorithm:

Aim: Voted for selected left-leaning parties

● Aim applies ● Aim does not apply ● Excluded by conditions

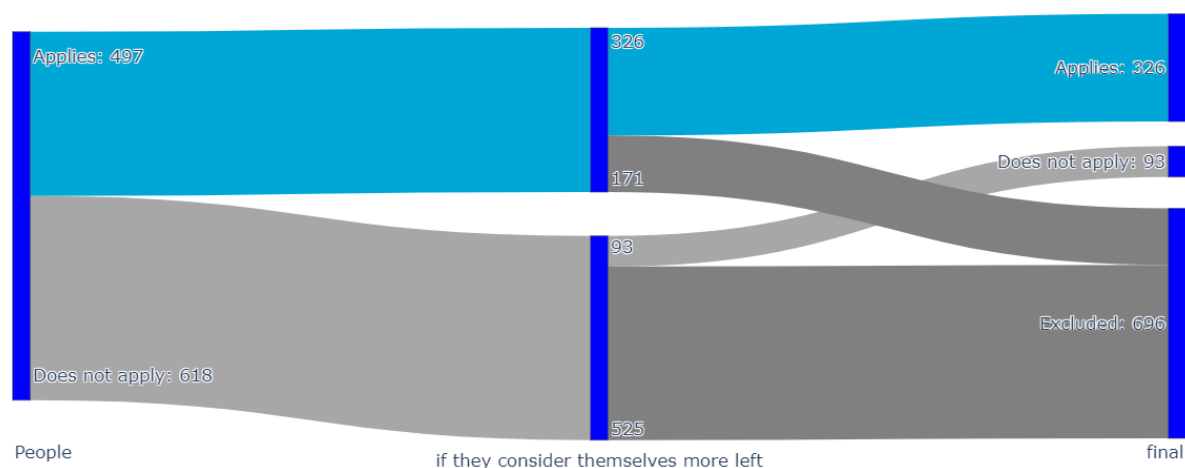


Figure 6.24: Sankey Diagram representing the first extracted informal rule from table 6.4 comparing voting behaviour with self placement on left-right political spectrum.

Figure 6.24 illustrates that out of 1115 total respondents, 497 voted for selected left-leaning parties. Among the 419 who consider themselves more left, 326 (77.8%) voted for left-leaning parties, while 93 did not. This shows a strong, but not absolute, correlation between left-leaning self-identification and voting for left-leaning parties.

Sankey plot representation of an IG statement extracted by decision tree algorithm:

Aim: Voted for other party

● Aim applies ● Aim does not apply ● Excluded by conditions

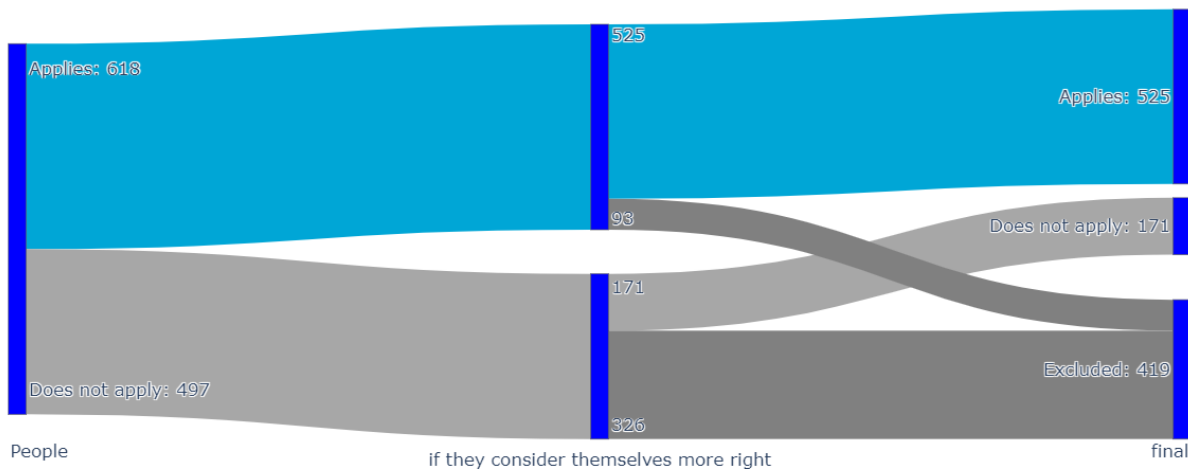


Figure 6.25: Sankey Diagram representing the second extracted informal rule from table 6.4 comparing voting behaviour with self placement on left-right political spectrum.

Figure 6.25 shows that out of 1115 respondents, 618 voted for other parties. Among the 696 who consider themselves more right, 525 (75.4%) voted for other parties, while 171 voted for left-leaning parties. This again demonstrates a strong, but not perfect, correlation between right-leaning self-identification and voting for non-left-leaning parties.

It's noteworthy that not all people who consider themselves left voted for the selected left-leaning parties. This could be because some left-leaning parties may not have been selected for the aim outcome grouping, and some of the selected parties could be considered centre-left, such as Democrats '66 (Otjes, S. 2018). Additionally, even though people consider themselves more left or right, their vote may deviate each election cycle due to various factors such as specific candidates, current events, or strategic voting.

Next, the decision tree algorithm is applied again where the sample size limit is reduced to 0.15 in order to see if more (weaker) informal rules will be extracted. An example decision tree from this last iteration is shown in Figure 6.26.

prtvthnl: Party voted for in last national election, Netherlands

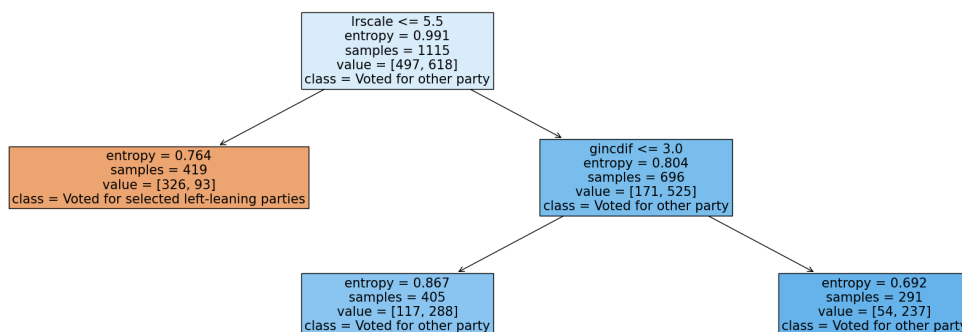


Figure 6.26: Example decision tree produced with the entropy threshold at 0.95 and sample size threshold at 0.15.

Figure 6.26 presents a decision tree with a depth of 2, similar to Figure 6.23 but with an added split on the right side. This additional depth allows for two conditions in the extracted informal rules. The tree's root node splits on the 'Irscale' variable (left-right scale), with 419 samples considering themselves more left and 696 samples more right. For those on the left, the majority (326) voted for selected left-leaning parties. The right branch further splits based on the 'gincdif' variable (government should reduce income differences), providing more nuanced insights into right-leaning voters' behaviour. The extracted informal rules from this final iteration are presented in Table 6.5.

Table 6.5: Table showing extracted informal rules when entropy threshold is 0.95 and sample size threshold set to 0.15.

Attribute	Aim	Condition 1	Condition 2	Entropy
People	Voted for other party	if they consider themselves more right	and if they do not agree that government should reduce differences in income levels	0.69
People	Voted for selected left-leaning parties	if they consider themselves more left	And none	0.76
People	Voted for other party	if they consider themselves more right	And none	0.80
People	Voted for other party	if they attend religious services	And none	0.81
People	Voted for other party	if they do not agree that government should reduce differences in income levels	And none	0.84
People	Voted for other party	if they consider themselves more right	and if they agree that government should reduce differences in income levels	0.87
People	Voted for selected left-leaning parties	if they rarely attend religious services	and if they think that key decisions are not made by national governments rather than the European Union	0.94

Table 6.5 presents seven extracted informal rules, showcasing a more diverse set of conditions influencing voting behaviour. Notably, the two informal rules from the previous Table 6.4 are also extracted here (rows 2 and 3), validating the consistency of the SIRE method across different threshold settings. The new rules provide additional insights:

1. The strongest rule (lowest entropy of 0.69) suggests that right-leaning individuals who disagree with government reducing income differences are highly likely to vote for other parties.
2. Religious service attendance emerges as a factor, with regular attendees more likely to vote for other parties (entropy 0.81).
3. Views on income inequality play a role, with those disagreeing with government intervention more likely to vote for other parties (entropy 0.84).
4. The weakest rule (highest entropy of 0.94) indicates that those who rarely attend religious services and prefer EU-level decision-making are more likely to vote for left-leaning parties.

These rules offer a more nuanced understanding of voting behaviour, incorporating multiple factors beyond just left-right political alignment. The newly extracted informal rules are represented in the Sankey plots of Figures 6.27 to 6.31.

Sankey plot representation of an IG statement extracted by decision tree algorithm:
Aim: Voted for other party

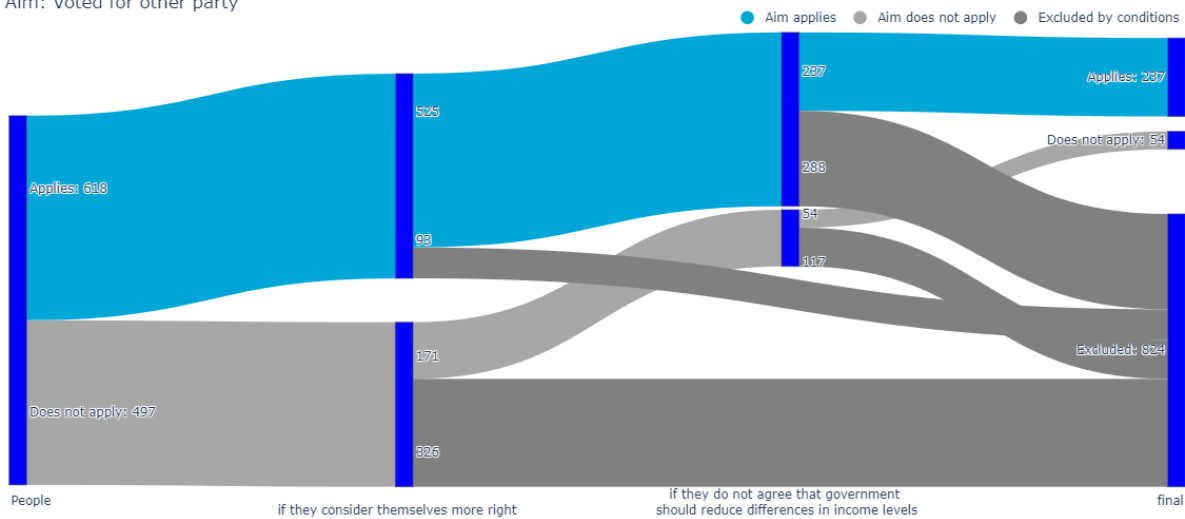


Figure 6.27: Sankey Diagram representing the first extracted informal rule from Table 6.5 comparing voting behaviour with self placement on left-right political spectrum.

Figure 6.27 illustrates the voting behaviour of those who consider themselves more right and disagree that the government should reduce income differences. Out of 1115 total respondents, 618 voted for other parties. Among the 696 who consider themselves more right, 291 also disagree with the government reducing income differences. Of these 291, a significant majority (237 or 81.3%) voted for other parties, demonstrating a strong correlation between right-leaning views, opposition to income redistribution, and voting for non-left parties.

Sankey plot representation of an IG statement extracted by decision tree algorithm:
Aim: Voted for other party

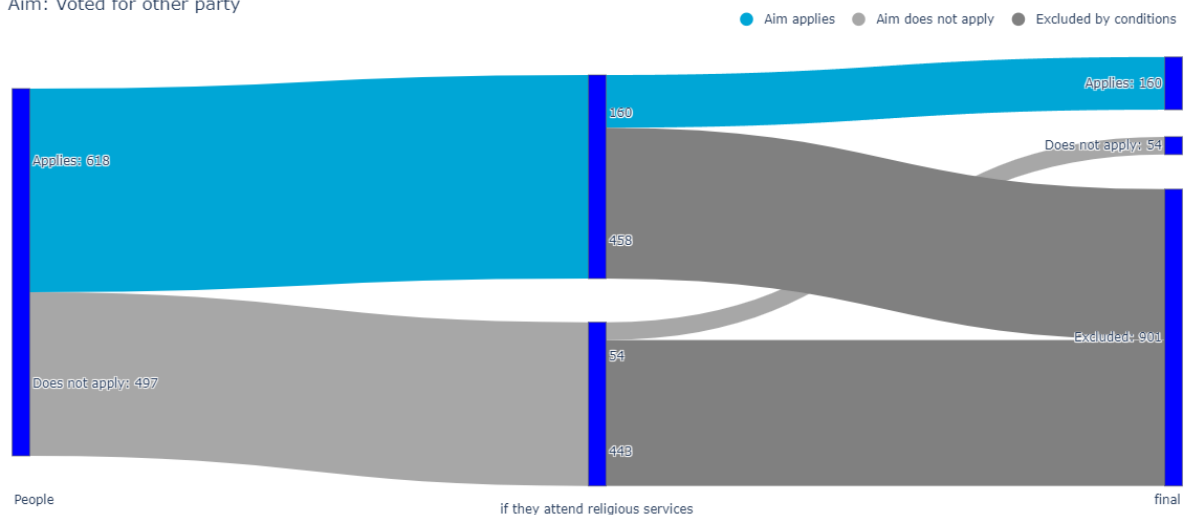


Figure 6.28: Sankey Diagram representing the fourth extracted informal rule from Table 6.5 comparing voting behaviour with self placement on left-right political spectrum.

Figure 6.28 shows the relationship between religious service attendance and voting behaviour. Out of 1115 respondents, 618 voted for other parties. Among the 214 who attend religious services, 160 (74.8%) voted for other parties. This indicates a tendency for regular religious service attendees to vote for non-left parties, although the relationship is not as strong as the left-right self-placement.

Sankey plot representation of an IG statement extracted by decision tree algorithm:
Aim: Voted for other party

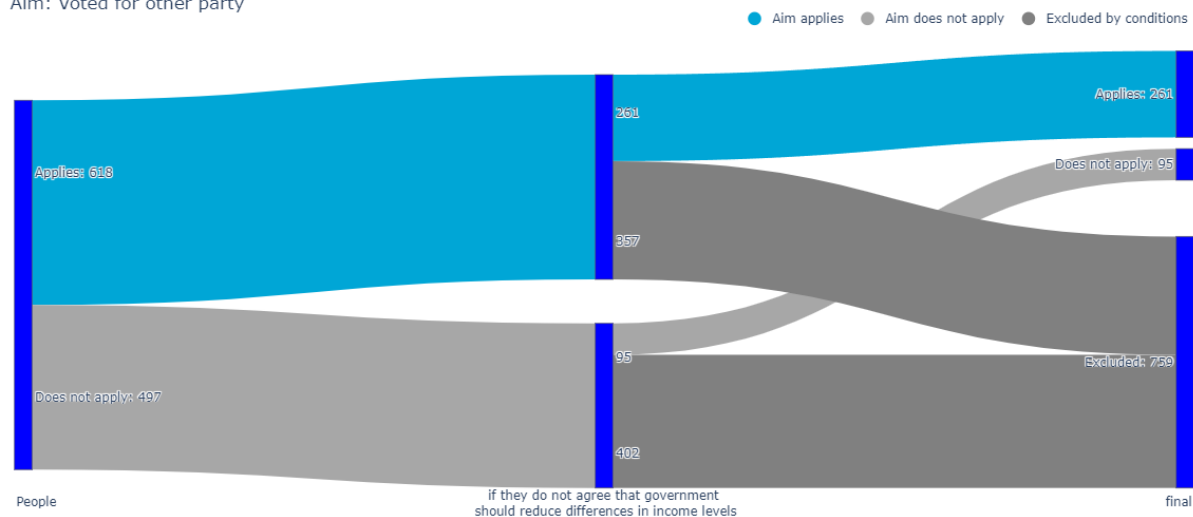


Figure 6.29: Sankey Diagram representing the fifth extracted informal rule from Table 6.5 comparing voting behaviour with self placement on left-right political spectrum.

Figure 6.29 illustrates the voting behaviour of those who do not agree that the government should reduce differences in income levels. Out of 1115 respondents, 618 voted for other parties. Among the 356 who disagree with the government reducing income differences, 261 (73.3%) voted for other parties. This shows a correlation between opposition to income redistribution and voting for non-left parties, though not as strong as the left-right self-placement.

Sankey plot representation of an IG statement extracted by decision tree algorithm:
Aim: Voted for other party

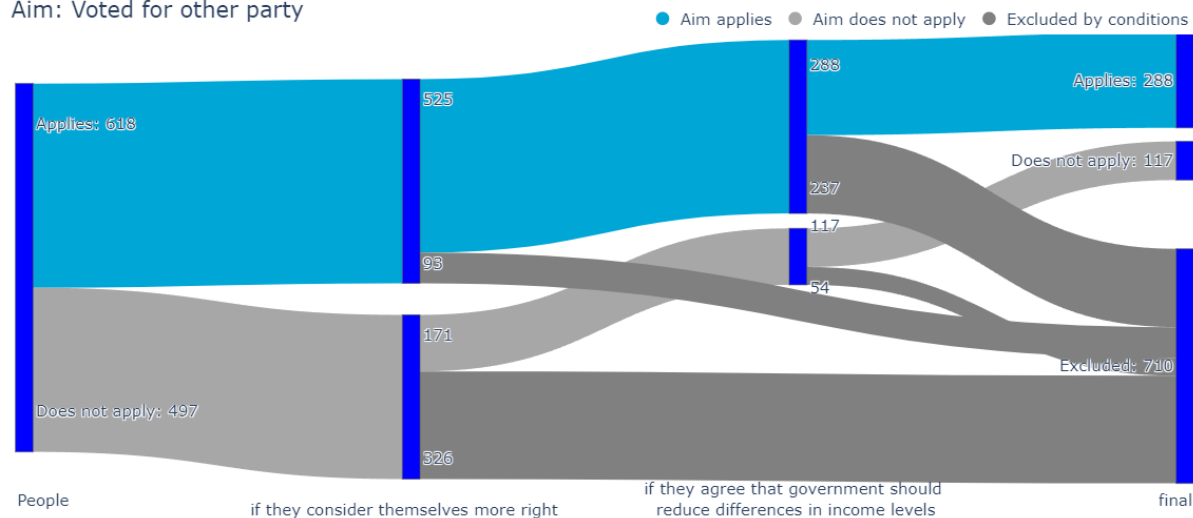


Figure 6.30: Sankey Diagram representing the sixth extracted informal rule from Table 6.5 comparing voting behaviour with self placement on left-right political spectrum.

Figure 6.30 depicts the voting behaviour of those who consider themselves more right and agree that the government should reduce differences in income levels. Out of 1115 respondents, 618 voted for other parties. Among the 696 who consider themselves more right, 405 also agree with the government reducing income differences. Of these 405, 288 (71.1%) voted for other parties. This rule shows that even among right-leaning voters who

support income redistribution, there's still a tendency to vote for non-left parties, although less pronounced than those who oppose redistribution.

Sankey plot representation of an IG statement extracted by decision tree algorithm:

Aim: Voted for selected left-leaning parties

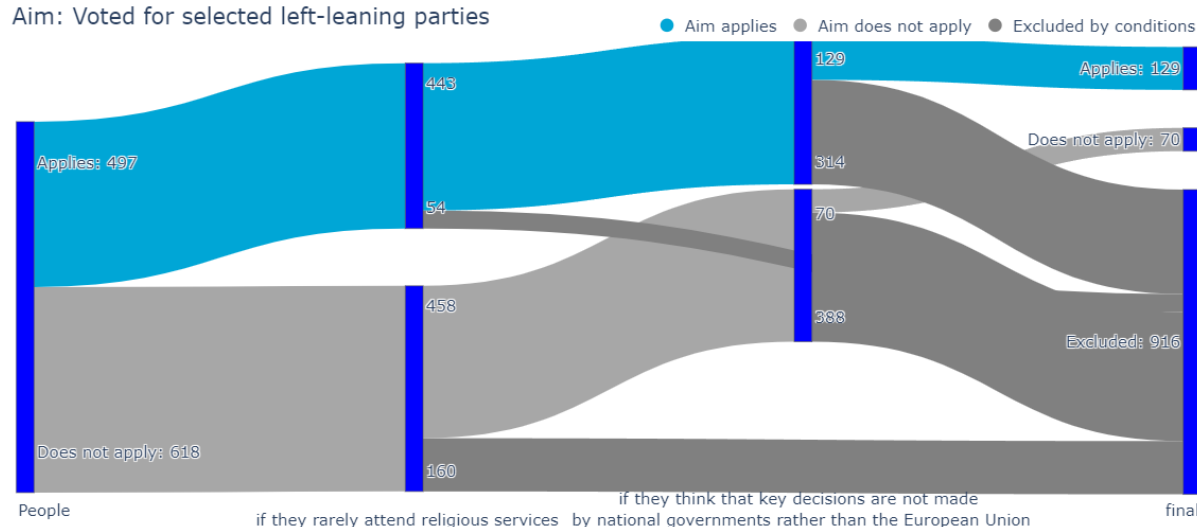


Figure 6.31: Sankey Diagram representing the seventh extracted informal rule from Table 6.5 comparing voting behaviour with self placement on left-right political spectrum.

Figure 6.31 shows the voting behaviour of those who rarely attend religious services and think that key decisions are not made by national governments rather than the European Union. Out of 1115 respondents, 497 voted for left-leaning parties. Among the 991 who rarely attend religious services, 199 also believe that key decisions are not made by national governments. Of these 199, 128 (64.3%) voted for left-leaning parties. This rule, while having the highest entropy, still shows a slight tendency for less religious, pro-EU individuals to vote for left-leaning parties.

The application of the SIRE method in this validation demonstrates its potential as a strategic tool for political analysis and policy development. Through a systematic approach of data preprocessing, decision tree analysis, and visualisation techniques, the method successfully identified key conditions linked to voting behaviour in the Netherlands. These insights offer valuable information that could inform policy or marketing strategies for political parties. The iterative process of refining the decision tree algorithm, adjusting sample size and entropy thresholds, and encoding survey responses proved effective in extracting meaningful informal rules.

Key strengths of the SIRE method revealed in this validation include its ability to handle complex, multi-faceted datasets like the ESS, flexibility in adjusting parameters to extract rules of varying strengths, visual representation of results through parallel set and Sankey diagrams, and capacity to uncover both expected and unexpected relationships in voting behaviour. However, the validation also highlighted areas for potential improvement, such as more sophisticated data preprocessing techniques, further refinement of encoding strategies, exploration of additional visualisation techniques, and better integration with domain-specific knowledge.

The SIRE method's application to voting behaviour analysis demonstrates its versatility beyond its original context of environmental management. Future applications in political science or other social science domains could benefit from collaboration with domain experts to refine question selection and interpretation, integration with other analytical techniques for more comprehensive analysis, longitudinal studies to track changes in informal rules over time and comparative studies across different countries or regions to identify broader patterns. With continued development and application, the SIRE method has the potential to become a robust and widely applicable approach for extracting informal rules from survey data, thereby informing more effective and targeted policy interventions in various social science domains.

6.4. Validation Summary

The SIRE method has demonstrated its validity and versatility across different domains, producing coherent and realistic results when compared with existing literature, hypothetical policy scenarios, and complex political data. The method successfully extracted logical informal rules reflecting common barriers and motivators in flood protection measures adoption and voting behaviour.

Key strengths of the SIRE method include:

1. Adaptability to diverse datasets, from specific behavioural surveys (SCALAR) to broad social attitudes (ESS).
2. Effective visualisation techniques (parallel set and Sankey diagrams) enhancing data interpretability.
3. Capacity to uncover both expected and unexpected relationships in complex social behaviours.
4. Flexibility in adjusting parameters to extract rules of varying strengths and significance.

The validation procedures confirmed that the SIRE method is robust and capable of producing meaningful insights for policy development and behavioural analysis across different contexts. Its ability to bridge qualitative and quantitative data analysis makes it a valuable tool for institutional analysts and policymakers. However, areas for improvement and further development include:

1. **Data Preprocessing:** Refine techniques for encoding and categorising variables to capture nuanced trends more accurately.
2. **Iterative Refinement Process:** Develop a more systematic approach for selecting and encoding relevant questions, especially for large, diverse datasets like the ESS.
3. **Integration with Domain Expertise:** Collaborate with subject matter experts to improve question selection and result interpretation.

In conclusion, while the SIRE method has proven effective, targeted improvements in data preprocessing, condition selection, and result interpretation can further enhance its applicability and effectiveness in extracting informal rules from survey data, thereby informing more nuanced and effective policy interventions in complex social-ecological systems.

7. Reflections on the SIRE Method

This chapter performs a comprehensive reflection on the Survey Informal Rule Extraction (SIRE) method. This discussion explores the practical applications of the SIRE method within IA research. Additionally, the significant role of IG in structuring and interpreting survey data is addressed. The chapter also highlights the potential of automating certain steps of the SIRE method using generative AI.

7.1. Use Cases in Institutional Analysis Research

The SIRE method generates output with numerous practical applications for Institutional Analysis (IA) research, offering valuable insights into behavioural patterns and their influencing conditions. This section discusses how these outputs can be utilised across various contexts, emphasising the distinction between policies and formal regulations.

The structured statements extracted by SIRE can be compared with existing policies to identify alignment or gaps. For example, in flood risk management, if a policy aims to increase community resilience, the extracted informal rules can highlight whether targeted behaviours are prevalent and under what conditions they occur. This comparison can inform policy refinement, helping policymakers understand which aspects of their policies are effective and which may need adjustment.

Importantly, not all policies translate directly into formal regulations. The SIRE method can help bridge this gap by:

1. **Identifying regulatory gaps:** Comparing extracted informal rules with existing regulations can reveal areas where formal rules are lacking or misaligned with public behaviour.
2. **Informing regulatory design:** Insights into factors influencing behaviour (e.g., age, perceived cost, self-efficacy) can guide the creation of more targeted and effective regulations.
3. **Assessing regulatory impact:** Evaluating how existing regulations align with or influence informal rules and behaviours provides insights into the effectiveness of current formal rules.

For instance, if SIRE reveals that perceived cost is a significant barrier to flood protection measure adoption, this could inform regulations mandating financial assistance programs or tax incentives.

The demographic proportions and behavioural rules extracted by SIRE can inform agent characteristics in agent-based models. For example, in modelling flood adaptation behaviours, agents can be programmed to mirror the decision-making patterns identified by SIRE, potentially enhancing the model's predictive power. This application is particularly relevant to studies like those by Wagenblast, T. (2022) and Lechner J. (2022), where more representative datasets and refined modelling approaches were needed. SIRE outputs can inform network structures and interactions in Institutional Network Analysis. For instance, in analysing institutional networks for climate change adaptation, the extracted informal rules

and their associated statistical values can help map the complexities of inter-organizational relationships and decision-making processes.

Decision tree values from SIRE highlight relationships between conditions and behaviours. As demonstrated in Chapter 6, researchers can identify which conditions most significantly influence specific behaviours, allowing for more targeted interventions. For example, if SIRE reveals that experiencing a flood increases the likelihood of purchasing flood insurance, policymakers might design interventions providing immediate post-flood support to encourage insurance uptake.

SIRE allows researchers to identify significant emergent behaviours by setting thresholds for population engagement. While this study initially used a 33% threshold inspired by Schelling's model, the validation process revealed that thresholds can be adjusted to discover weaker relationships. For instance, in the SCALAR dataset analysis (Section 6.1), a 25% threshold yielded comprehensive results while maintaining statistical significance. This flexibility enables researchers to calibrate the method to their specific context.

Visualisation tools like parallel sets and Sankey diagrams are crucial for interpreting and communicating SIRE results. As shown in Chapters 5 and 6, these visualisations can reveal complex relationships between conditions and behaviours, aiding in understanding the cumulative impact of multiple policies or interventions. For example, in the ESS data analysis (Section 6.3), visualisations highlighted significant demographic and attitudinal factors influencing voting patterns, providing actionable insights for political strategy.

By understanding conditions leading to specific actions, SIRE can help predict population responses to events or policies. This is valuable for designing proactive and adaptive measures. For instance, in flood management, SIRE can predict the proportion of a demographic likely to engage in protective behaviours following a flood event, enabling more effective resource allocation and intervention design (Abebe, Y. A. et al., 2020). It is important to note that predictions cannot be easily made with cross-sectional data. Instead, longitudinal data, which involves repeated observations over time, allowing for the analysis of changes and trends, should be used (Longford, N. T. 2008). For example, this could be achieved by performing the SIRE method on multiple waves of the SCALAR dataset.

In conclusion, the SIRE method's outputs offer a robust foundation for advancing IA research across various domains. By integrating these outputs into policy advice, agent-based models, and institutional design, researchers can gain richer insights into social behaviours and their influencing conditions, leading to more informed and effective policy and institutional designs. The method's ability to bridge the gap between informal rules and formal regulations makes it particularly valuable for policymakers and institutional analysts seeking to create more effective and responsive governance structures.

7.2. The Role of Institutional Grammar

IA is a methodological approach that examines the rules, norms, and strategies governing behaviour and interactions within institutions. IG plays a pivotal role in IA by providing a structured syntax to decompose and analyse these institutional components. IG facilitates

the systematic examination of institutional statements, categorising them into attributes, deontics, aims, conditions, and or else components. These components offer a clear understanding of how institutions operate and influence behaviour (Frantz, C. et al. 2021).

A key topic of this thesis was understanding how different types of survey questions can capture IG components. Based on findings from previous sections and the review of existing surveys, different types of questions were evaluated for their effectiveness in extracting information about institutional informal rules. The types of values include dichotomous, categorical, ordinal, open numerical, or open text-based, each requiring different statistical analyses, visualisation approaches, and methods to link them to ADICO components. The complexity of response options influences the difficulty of data analysis, as represented in figure 7.1.

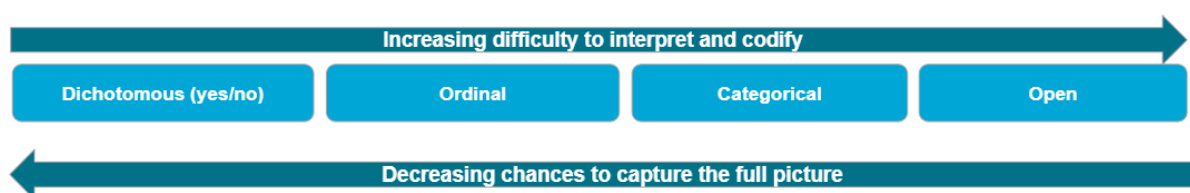


Figure 7.1: Visualisation of the difficulty to codify trend of question types.

Due to the difficulty to interpret and encode certain questions, the use of IG led to a preference for dichotomous and ordinal questions for the developed method. Categorical questions can be implemented in the SIRE method but open questions are ignored entirely. On that note, open questions overlap in content and structure with interview transcripts and require a similar analysis method which is another reason why they were excluded from the scope of this thesis.

Potential limitations to using IG have been listed in Appendix H. While it is important to note that there are limitations to using Institutional Grammar, the goal of this thesis was to create a method to link Institutional Grammar with survey data. This work contributes to the existing body of knowledge by providing researchers with a novel tool and approach. If researchers are unsure if this is a suitable approach for their research, they can refer to Appendix E.

If IG had not been utilised in this study, the outcome would likely have been significantly different. Without IG, the analysis of informal rules would have lacked a standardised framework, potentially leading to a less coherent and more fragmented understanding of the data. Other methodologies could have been employed to identify and interpret informal rules. For instance, agent-based modelling focuses on the characteristics and behaviours of individual agents within a system, offering insights into emergent patterns and interactions. General statistical analyses and descriptive methods might identify trends and correlations within survey data, providing a broad overview of the data's characteristics. INA examines network properties and relationships within institutional settings, highlighting connections and influences among different actors and entities.

Avoiding IG entirely would have resulted in an analysis that might miss the nuanced relationships between various components of institutional rules. The absence of a structured grammar would make it challenging to compare findings across different datasets or to generalise the results to broader contexts. This lack of standardisation could lead to difficulties in replicating the study or applying its findings to real-world policy-making and IA.

Furthermore, the use of a standardised framework such as IG ensures that the results are transferable to alternative approaches, which may not have been as feasible in reverse (e.g., transferring agent-based model characteristics to IG).

This thesis ultimately used IG to offer a reproducible framework for extracting and analysing informal rules from survey data. IG's structured approach ensures that the analysis is not only systematic but also scalable, making it possible to apply the method to large datasets consistently. Additionally, IG's ability to delineate the details of institutional statements allows for a clear understanding of the underlying social structures and behaviours. This depth of analysis is crucial for informing evidence-based policy-making and designing effective interventions, goals that are central to the objectives of this thesis.

7.3. Automating Steps with Generative Artificial Intelligence

Generative AI has experienced a transformative leap in the capabilities available to researchers, particularly in fields involving extensive text processing and generation. This technology has the potential to make tasks such as text processing faster. However, it is crucial to recognize the importance of transparency and methodological rigour when integrating AI into research processes. Researchers must be well-informed about these new tools and explicitly document their use in their work to maintain the integrity and reproducibility of scientific research.

The automation of question categorization proved to be both challenging and, in many cases, superfluous. In practice, researchers often manually select which questions to analyse and designate as aims, with the remaining questions typically categorised as conditions. While AI techniques, such as text analysis and machine learning models, have the potential to automate this categorization process, our experiments with these approaches yielded inconsistent results. The accuracy of the AI-based categorization varied significantly and required substantial refinement to be reliable. Given these limitations and the fact that manual selection often suffices, a decision was made not to include an automated categorization step in the final design of the SIRE method. This decision prioritises researcher discretion and contextual understanding over potentially unreliable automation, ensuring the integrity of the initial data organisation process.

The conversion of survey questions into ADICO syntax is another area where AI can provide significant benefits. This involves generating structured statements that clearly define the institutional components. While this process has not yet been perfected, the approach maintains consistency among the statements and requires separate statements to be made for each split of ordinal responses (e.g., age ranges). AI can automate this rephrasing process, ensuring that the conversion is both accurate and efficient (Törnberg, P. 2023). Converting the text of survey questions into IG statement components is a largely deterministic and objective process. Since the goal is to rephrase the questions into a specific format without adding new information, the risk of introducing bias is minimised. This deterministic nature ensures that the conversion process remains consistent and reliable, providing clear definitions of IG components with little room for variation (Ziems, C. et al., 2024).

AI can also be employed to identify correlations between conditions and aims within survey responses. AI can uncover relationships and dependencies that might be difficult to detect manually. However, this step involves the greatest risk of "black box" confusion, where the internal workings of AI models are not transparent. Careful validation and interpretation of the results are essential to mitigate this risk. The importance of human supervision and task-specific prompt engineering to ensure the reliability of AI outputs in social science research remains (Bail, C. A. 2024). The current implementation with decision trees and parallel set diagrams is a clear approach with useful outputs, therefore reducing the necessity to include any form of AI at this stage.

In conclusion, the integration of generative into IA presents a promising avenue for enhancing research efficiency and accuracy. By automating the selection, categorization, conversion, and relationship identification processes, researchers can focus more on the interpretation and application of findings. When integrating AI into these processes, it is essential to acknowledge and address the associated challenges and limitations. For instance, bias in training data, ethical concerns, and the potential for low-quality research outputs are significant issues that need careful consideration. However, it is crucial to use these tools transparently and responsibly, ensuring that the methodologies developed facilitate the effective use of modern tools while maintaining the rigour and reproducibility of scientific research. Creating open-source infrastructure for research on human behaviour can help mitigate these challenges by ensuring broad access to high-quality research tools and fostering a deeper understanding of the social forces that guide human behaviour (Bail, C. A. 2024).

8. Conclusions

This final chapter synthesises the findings of this thesis, concluding the effectiveness and limitations of the developed method. Additionally, it outlines future research directions and areas for further exploration, laying the groundwork for advancing knowledge in IA methodologies. First, the research questions will be addressed.

8.1 Research Questions

This section reflects on the research questions posed developed for this thesis, evaluating their answers based on the findings and methodologies applied throughout the study. This thesis' main research question is:

What standard method can be developed to extract informal rules in use for institutional analysis from survey data?

By answering this question and its associated sub-questions, this thesis has introduced a novel methodology named the SIRE (Survey Informal Rules Extraction) method. This method offers a structured approach for future researchers, particularly institutional and policy analysts, to extract and analyse behavioural data systematically from surveys.

SQ1. ***What challenges can institutional and policy analysts face when identifying institutional informal rules?***

The first sub-question (SQ1) examines the challenges institutional analysts face when identifying institutional informal rules. The research revealed that data complexity, subjectivity, and scalability are major challenges in identifying informal rules. The SIRE method addresses these by providing a structured, reproducible approach that integrates quantitative data analysis, potentially enhancing the rigour and scalability of IA. Through an extensive literature review and gap analysis, several critical challenges emerged. To address SQ1, a literature review and gap analysis were conducted. The review highlighted several challenges faced by institutional analysts:

1. **Data Complexity and Diversity:** Analysts often struggle with the complexity and diversity of socio-technical systems, which require robust methods to accurately capture informal rules.
2. **Subjectivity and Bias:** Existing methods heavily rely on qualitative data, which can introduce subjectivity and bias.
3. **Scalability:** Traditional approaches are labour-intensive and not easily scalable, limiting their application to larger datasets.

The gap analysis further identified underexplored areas, such as the integration of quantitative data and the need for standardised methods. The SIRE method developed in this study addresses these limitations by providing a structured, reproducible approach for extracting informal rules from survey data. By integrating quantitative data and leveraging statistical analysis techniques, the SIRE method has the potential to enhance the rigour and scalability of IA, mitigating the identified challenges.

SQ2. *To what extent does Institutional Grammar align with survey data to identify and organise informal rules?*

This thesis demonstrates that Institutional Grammar (IG) aligns well with survey data for identifying and organising informal rules. The structured approach of IG, particularly its ADICO components (Attributes, Deontics, Aims, Conditions, Or else), provides a robust framework for organising institutional content derived from survey responses. IG aligns well with survey data, providing a robust framework for organising institutional content. The research discovered that survey questions can be effectively mapped onto ADICO components, though some components (like deontics and or else) present challenges in direct survey application. This alignment opens new possibilities for systematic analysis of informal rules in survey data.

In the context of IG and the method developed in this research, the primary goal of a survey is to identify the various components of ADICO. Aggregating the responses results in statistically linked questions and components which can be translated into informal rules. To understand the structure and questioning used in surveys, three surveys that are frequently used for sociological and institutional research were examined: SCALAR, the ESS, and the Krefeld-Schwalb et al. (2024) survey data. A key reason for conducting surveys is to understand a demographic's preferences, behaviours, and opinions. These elements are directly reflected in the ADICO informal rule structure. A pattern observed from reviewing existing surveys is that questions tend to be ordered according to these ADICO elements:

Table 8.1: Categorization of questions, their patterns and roles in survey design and analysis.

Attributes:	Questions to distinguish relevant demographics.	Who? e.g., "people"
Aims:	Questions about actions performed, intended actions, changes made, or opinions held.	What? e.g., "purchase an electric vehicle"
Conditions:	Questions about relevant factors impacting responders' behaviour.	When or why? e.g., "if their neighbour has one"
Deontics:	Questions about social obligations or permissions.	Duty? e.g., "may"
Or else:	Uncommon questions, but there could be questions about policy enforcement of behaviour.	And if not? e.g., "or else..."

As shown in Table 8.1, survey questions can be effectively mapped onto IG components:

- Attributes are typically identified through demographic questions (age, location, gender, etc.), which are essential for examining policies and behaviours applicable to specific groups. These questions can be categorical or ordinal.
- Aims are captured through questions about past actions, future intentions, or current beliefs. These directly ask whether respondents have taken specific actions, intend to take those actions, or hold particular opinions.
- Conditions are identified through questions about external factors that could influence actions and beliefs. These include experiences, financial matters, and social conditions that play a significant role in shaping behaviour. Questions often ask if certain conditions influenced behaviour regarding a specific aim.
- Deontics, while challenging to capture directly, can be inferred from the language used in questions. Key phrases like "are you allowed to," "do you have the ability to," and "are

you expected to" can indicate deontics. However, implementing these elements is currently beyond the scope of this thesis.

- Or else components are typically extracted from policy documents as regulatory statements. However, survey questions can prescriptively ask about whether different "or else" scenarios would influence respondents' actions, using dichotomous or Likert scale questions to assess the impact of fines, subsidies, or access revocation on behaviour.

The practical application of IG to survey data was demonstrated through the analysis of existing datasets. The results indicated that IG could be successfully applied to structure survey responses and extract meaningful informal rules. The combined use of qualitative data analysis and statistical methods further validated the applicability of IG and enhanced the extraction and organisation of informal rules from survey data.

SQ3. *What are the possible approaches to connect survey results to Institutional Grammar concepts?*

The third sub-question (SQ3) investigates the possible approaches to connect survey results to IG concepts. The study identified several effective methodologies through a blend of literature review and practical testing. Key to this process was understanding how different types of survey questions can capture IG components. Based on findings from previous sections and the review of existing surveys, different types of questions have been evaluated for their effectiveness in extracting information about institutional informal rules.

- **Dichotomous:** Clear, straightforward responses; ideal for identifying demographics, actions, or conditions but may exclude unasked factors.
- **Ordinal:** Measures opinions or factors in order; useful for identifying conditions; can be numerically encoded and analysed using decision trees.
- **Categorical:** Multiple predefined response options; reveals a range of insights into conditions or attributes; each response needs separate handling.
- **Open:** Detailed responses; informs conditions and aims; requires language processing to cluster general outcomes; difficult to obtain concrete results.

Open questions offer the advantage of discovering diverse ideas and can lead to outcomes that the researcher might not have initially considered. They are valuable for uncovering additional information and contributing to a deeper analysis of social behaviours and institutional contexts. However, for IA, where the focus is on understanding the impacts of specific policies and factors, more structured and easier-to-codify question types are often more effective. These structured questions provide clearer, more straightforward data that can be efficiently processed and analysed to extract ADICO components.

In conclusion, understanding the types of survey questions and their effectiveness in extracting ADICO components is crucial for developing a method to map survey data to IG components. While open questions can enrich the analysis with unexpected insights, prioritising easier-to-codify question types ensures a more streamlined and effective extraction process. The survey design emerged as a critical factor in accurately capturing IG components. Practical tests on existing surveys revealed that well-crafted questions are essential for aligning survey data with IG.

SQ4. *What requirements should the output data fulfil to ensure that the informal rules from surveys are usable for Institutional Analysis, consistent and well-communicated?*

The research established that output data must be clear, practical, and contextualised within the broader dataset. Visualisation techniques like Parallel Set Plots and Sankey diagrams proved effective in representing complex relationships. The iterative development process revealed the importance of continuous refinement and validation to ensure reliability and alignment with theoretical constructs.

Initially, a review of existing data analysis techniques and feature extraction methods informed the development of a suitable extraction method tailored to the dataset. The development and refinement process incorporated both qualitative and quantitative techniques, ensuring continuous improvement through rigorous testing and validation. Then the effectiveness of the method was evaluated by comparing the extracted structures to known institutional data and literature, confirming the reliability and alignment of the outputs with theoretical constructs.

The extracted informal rule statements should indicate the proportion of responders it represents, as this contextualises the rule within the broader dataset. The data should support various use cases, including agent-based models, network diagrams, and policy evaluation. The SIRE method ensures that each of these applications is provided with the necessary information to enhance results.

Parallel Set Plots plots effectively represent the initial distribution of responses, making it easy to interpret relationships and flows between different questions. Decision Trees are Generated using a pre-existing Python package-based algorithm, these need further refinement to improve readability and interpretability. To further improve the SIRE method a custom-designed algorithm be developed specifically for the SIRE method to optimise visualisation for IA. Sankey diagrams clarify selected Institutional Grammar (IG) statements and the progression in sample size and proportion, ensuring the data is presented clearly and simply.

Overall, this research demonstrates the feasibility and value of integrating IG with survey data analysis, offering a novel approach to understanding complex socio-technical systems through the lens of informal rules. The SIRE method provides a scalable, reproducible framework that addresses key challenges in IA, opening new avenues for research and policy analysis.

8.2. Contributions

This thesis makes significant contributions to the field of Institutional Analysis (IA) by developing the SIRE (Survey Informal Rules Extraction) method, a standardised approach for extracting informal rules and social structures from survey data. The SIRE method addresses several key challenges in IA, offering a new procedure that can be consistently applied to survey data, facilitating the identification and formalisation of informal rules. This innovative methodology not only answers the primary research question but also provides a

comprehensive guide for extracting actionable insights from survey responses, thus enriching the methodological toolkit available for IA.

A key strength of the SIRE method lies in its integration of quantitative survey data with qualitative Institutional Grammar concepts. This unique synthesis bridges a critical gap in the field, allowing for a more nuanced understanding of social structures and behaviours. By providing a systematic way to analyse large datasets, SIRE enables researchers to conduct more comprehensive studies of institutional dynamics, potentially uncovering patterns and relationships that might be missed by traditional qualitative methods.

The quality and validity of the results produced by the SIRE method have been thoroughly evaluated through its application to existing survey datasets such as SCALAR and ESS. This cross-domain validation demonstrates the method's versatility and robustness, enhancing its credibility and potential for widespread adoption. The validation process, which included testing against case studies, ensures that the results are not only informative but also robust and reliable. The outcomes obtained are indicative of underlying social structures and informal rules, offering valuable insights for further research and policy-making.

One of the most significant contributions of this thesis is the reproducibility of the SIRE method. Detailed instructions for its application are provided in Chapter 5, along with a Python script to facilitate implementation. This level of detail ensures that other researchers can replicate the study and achieve similar results, establishing SIRE as a robust framework that can be widely adopted in the field of IA. This reproducibility is crucial for advancing the field, as it allows for consistent comparison of results across different contexts and datasets.

The introduction of the SIRE method has significant implications for future research in IA and related fields. By providing a standardised and streamlined approach, SIRE acts as a catalyst for further studies, enhancing the efficiency and comparability of research outcomes. Researchers can now design their questionnaires with this method in mind, potentially leading to improved data quality and more insightful analyses of social structures. Furthermore, the method's framework for survey design offers guidelines that could improve the quality and relevance of data collection in IA.

The practical applications of the SIRE method are diverse and impactful. It facilitates the extraction of structured behavioural information from survey data, which can be directly utilised in complementary research steps such as policy advice, the design of agent-based models, and other institutional representations. By formalising informal rules, the method provides a practical tool for researchers and policymakers to understand and influence social behaviours more effectively. This bridge between academic research and practical policy implementation is a significant contribution to the field.

In conclusion, this thesis contributes a significant and innovative methodology to the field of IA. The SIRE method enhances current research practices, offering a reliable, reproducible, and efficient means of extracting informal rules from survey data. By doing so, it supports the development of more nuanced and effective institutional analyses, ultimately contributing to better-informed policy-making and a deeper understanding of social structures. The potential for SIRE's application in related fields such as sociology, political science, and public policy further underscores its value as a versatile and powerful tool for social science research.

8.3. Limitations

While the SIRE method represents a significant advancement in extracting informal rules from survey data, it is important to acknowledge and address its limitations. These constraints not only provide a balanced view of the method's capabilities but also highlight areas for future refinement and research.

A fundamental limitation lies in the method's current approach to defining informal rules. The threshold for determining significant behavioural patterns based on the number of respondents meeting certain conditions requires further fine-tuning. This aspect of the method introduces a degree of subjectivity that could impact the consistency of results across different studies. Researchers implementing SIRE are encouraged to adjust these thresholds to suit their specific research contexts, which, while offering flexibility, may also lead to challenges in comparing results across studies. Future work should focus on establishing more robust guidelines for setting these thresholds, perhaps through meta-analyses of multiple SIRE applications.

The method's validation, primarily conducted using two existing datasets (SCALAR and ESS), potentially limits its generalizability. These datasets, while comprehensive, represent specific types of questions and survey structures. As SIRE is applied to different contexts and survey designs, new limitations may emerge, necessitating further adaptation and testing. To fully assure the method's reliability and robustness, it is crucial to evaluate and refine SIRE across a broader range of datasets and scenarios, including those from diverse cultural and socio-economic contexts.

The integration of AI in the SIRE method, while enhancing efficiency in text processing and rewriting, introduces its own set of challenges. The current implementation may lead to inconsistencies, particularly in text-processing tasks, and could potentially obscure the method's transparency. The reliance on AI in research methodologies requires careful consideration of potential biases and the need for human oversight (Hajikhani, A., et al. 2024). The process of rewriting survey questions and responses into ADICO components, particularly for ordinal questions, needs further refinement to enhance flexibility and accuracy.

The use of decision trees in identifying influential conditions, while a strength of the method, also presents limitations. The current approach may overlook alternative combinations of conditions that could provide valuable insights. The iterative process of repeating the decision tree analysis with previously selected conditions removed is a temporary solution that requires further development. Future versions of SIRE should explore more sophisticated algorithms or machine learning techniques that can capture a wider range of relevant conditions without compromising the method's interpretability.

A significant constraint of the current SIRE method is its limited ability to handle open-ended questions and identify deontic components. Open-ended responses, which often provide rich qualitative data, currently require separate analysis outside the SIRE framework. Similarly, the method's current inability to systematically identify deontic elements (specifying obligations or permissions) limits its comprehensiveness in capturing the full spectrum of

institutional rules. Addressing these limitations would significantly enhance the method's applicability across various survey types and institutional contexts.

The technical implementation of SIRE, currently tailored to specific survey formats and requiring Python coding skills, may pose accessibility challenges for some researchers. This limitation could restrict the method's broader adoption, particularly among researchers less familiar with programming or data preprocessing techniques. Developing a more user-friendly interface or providing more flexible data input options could significantly enhance the method's accessibility and adoption.

In conclusion, while the SIRE method offers a valuable new approach to extracting informal rules from survey data, addressing these limitations is crucial for its continued development and broader application. Future research should focus on refining the method's core algorithms, expanding its capabilities to handle diverse question types, improving its AI integration, and enhancing its user accessibility. By addressing these constraints, SIRE can evolve into an even more robust and widely applicable tool for institutional analysis, contributing significantly to our understanding of complex social systems and informing evidence-based policymaking.

8.4. Future work

The SIRE method, in its initial iteration, presents significant opportunities for refinement and expansion. Future research should focus on several key areas to enhance its robustness, applicability, and effectiveness in institutional analysis.

A primary focus for future work should be the development of a custom decision tree algorithm specifically designed for SIRE. This tailored algorithm would better accommodate categorical questions, which are common in survey data but can be challenging for standard decision tree algorithms. The custom algorithm could incorporate methods to handle response splits more effectively.

To ensure the SIRE method's reliability and generalizability, it is crucial to test it across a diverse array of survey datasets. This process will help identify compatibility issues and specific limitations, guiding necessary adjustments. By expanding the range of datasets, the method can be fine-tuned to handle various survey formats and question types effectively, enhancing its versatility and applicability. This extensive validation will contribute to the method's overall robustness and its ability to provide meaningful insights across different contexts.

Future research should also explore the SIRE method's potential in synthesising policy advice. By assessing current community interactions with existing policies, the method can provide valuable insights for policymakers. This application will contribute to the design of more effective and targeted policy interventions, bridging the gap between academic research and practical policy implementation. The informal rules extracted through the SIRE method offer a novel input for agent-based models in defining agent characteristics. Studies should focus on integrating SIRE outputs into the design of these models, enhancing the accuracy of simulations of complex social systems.

Exploring the potential of incorporating SIRE method results into other institutional analysis techniques, such as Institutional Network Analysis (INA), could provide a more comprehensive understanding of institutional interactions. This integration could lead to more nuanced and detailed representations of institutional landscapes. Additionally, future work should focus on developing precise techniques to extract deontics from survey data, enabling better identification of norms. This advancement will make the method more robust and comprehensive, providing a fuller picture of institutional structures and behaviours.

Integrating open-ended questions more effectively into the SIRE method is another crucial area for development. This would allow for a richer extraction of informal rules and exploration of underlying objectives and interests of actors, providing deeper insights into the institutional landscape. As AI technology evolves, future studies should explore the integration of advanced tools, such as large language models, to automate and enhance the extraction of informal rules. While embracing these advancements, it's crucial to maintain transparency and interpretability in the analysis process.

Developing standardised survey structures and question formats specifically aligned with the SIRE method is a promising area for future research. This standardisation will facilitate consistent and accurate extraction of informal rules, improving the quality and comparability of research outcomes. By designing surveys tailored to institutional analysis, researchers can enhance the efficacy of the SIRE method and its applicability across various fields of study.

In conclusion, the SIRE method presents a solid foundation for advancing institutional analysis through quantitative survey data. By pursuing these avenues of future research, particularly the development of a custom decision tree algorithm and the expansion of its capabilities to handle diverse question types, the method can be refined and expanded, offering increasingly valuable insights into social behaviours, institutional structures, and policy effectiveness. This ongoing development will contribute significantly to the field of Institutional Analysis, bridging theoretical frameworks with practical applications in policy-making and social system modelling. As the method evolves, it has the potential to become an indispensable tool for researchers, policymakers, and practitioners seeking to understand and shape institutional dynamics in complex social systems. The SIRE method not only advances our current understanding of informal rules but also opens up new possibilities for more nuanced, data-driven approaches to institutional analysis, promising a future where policy decisions are informed by a deeper, more comprehensive understanding of social structures and behaviours.

References

- Abebe, Y. A., Ghorbani, A., Nikolic, I., Manojlovic, N., Gruhn, A., & Vojinovic, Z. (2020). The role of household adaptation measures in reducing vulnerability to flooding: a coupled agent-based and flood modelling approach. *Hydrology and Earth System Sciences*, 24(11), 5329–5354. <https://doi.org/10.5194/hess-24-5329-2020>
- Backhouse, J. K. (1984). Inferential and non-inferential tests, *Educational Research*, 26:1, 52-55, <https://doi.org/10.1080/0013188840260109>
- Bail, C. A. (2024). Can Generative AI improve social science? *Proceedings of the National Academy of Sciences*, 121(21), e2314021121. <https://doi.org/10.1073/pnas.2314021121>
- Bender, E. M., Gebru, T., McMillan-Major, A., & Shmitchell, S. (2021). On the Dangers of Stochastic Parrots: Can Language Models Be Too Big? *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency*, 610–623. <https://doi.org/10.1145/3442188.3445922>
- Bognøy, J. (2021). IG Coder: Enabling Visual Coding of Institutional Statements (Master's thesis). Norwegian University of Science and Technology, Department of Computer Science. <https://hdl.handle.net/11250/2778079>
- Breza, E., Chandrasekhar A. G. & Larreguy, H., (2014). Social Structure and Institutional Design: Evidence from a Lab Experiment in India. <https://www.nber.org/papers/w20309>
- Chan, F. K. S., Yang, L. E., Mitchell, G., Wright, N., Guan, M., Lu, X., Wang, Z., Montz, B., & Adekola, O. (2022). Comparison of sustainable flood risk management by four countries - the United Kingdom, the Netherlands, the United States, and Japan - and the implications for Asian coastal megacities. *Natural Hazards and Earth System Sciences*, 22(8), 2567-2588. <https://doi.org/10.5194/nhess-22-2567-2022>
- Chueri, J., & Damerow, A. (2023). Closing the gap: How descriptive and substantive representation affect women's vote for populist radical right parties. *West European Politics*, 46(5), 928–946. <https://doi.org/10.1080/01402382.2022.2113219>
- Crawford, S. E. S. & Ostrom, E. (1995): A Grammar of Institutions, *Am. Polit. Sci. Rev.*, 89, 582–600, <https://doi.org/10.2307/2082975>.
- Eberly College of Science. (2023). STAT 504: Analysis of Discrete Data - 6.1 Introduction to GLMs. The Pennsylvania State University. <https://online.stat.psu.edu/stat504/lesson/6/6.1>
- Erdlenbruch, K., & Bonté, B. (2018). Simulating the dynamics of individual adaptation to floods. *Environmental Science Policy*, 84, 134-148. <https://doi.org/10.1016/j.envsci.2018.03.005>
- European Social Survey. (2024). Methodology Overview | European Social Survey. <https://www.europeansocialsurvey.org/methodology/methodology-overview>

Frantz, C., & Siddiki, S. (2022). Institutional Grammar: Foundations and Applications for Institutional Analysis. Palgrave Macmillan. <https://doi.org/10.1007/978-3-030-86372-2>

Gandhi, R., Saini, A., & Gaikwad, S. (2024). A Framework for Abstractive Text Summarization Using Hugging Face Transformers. 2024 14th International Conference on Cloud Computing, Data Science & Engineering (Confluence), 690–695. <https://doi.org/10.1109/Confluence60223.2024.10463423>

Gil-Clavel, S., Wagenblast, T., & Filatova, T., (2023). "Farmers' Incremental and Transformational Climate Change Adaptation in Different Regions: A Natural Language Processing Comparative Literature Review," SocArXiv 3dp5e, Center for Open Science. <https://doi.org/10.31235/osf.io/3dp5e>

Gittins, G. (2024). GeorgeGitHubbins/MScThesis-Extracting-InformalRules-from-Surveys. <https://github.com/GeorgeGitHubbins/MScThesis-Extracting-InformalRules-from-Surveys>

Grill, C. (2017). Longitudinal Data Analysis, Panel Data Analysis. In The International Encyclopedia of Communication Research Methods. <https://doi.org/10.1002/9781118901731.iecrm0134>

Grimm, V., Berger, U., Bastiansen, F., Eliassen, S., Ginot, V., Giske, J., DeAngelis, D. L. (2006). A standard protocol for describing individual-based and agent-based models. Ecological Modelling, 198(1-2), 115–126. <https://doi.org/10.1016/j.ecolmodel.2006.04.023>

Haer, T., Botzen, W. W., & Aerts, J. C. (2016). The effectiveness of flood risk communication strategies and the influence of social networks—insights from an agent-based model. Environmental Science Policy, 60, 44-52. <https://doi.org/10.1016/j.envsci.2016.03.006>

Hajikhani, A., & Cole, C. (2024). A Critical Review of Large Language Models: Sensitivity, Bias, and the Path Toward Specialized AI. Quantitative Science Studies, 1–22. https://doi.org/10.1162/qss_a_00310

IGRI. (2023). Institutional Grammar. Institutional Grammar Research Initiative. <https://institutionalgrammar.org/>

James, G., Witten, D., Hastie, T., Tibshirani, R., & Taylor, J. (2023). An Introduction to Statistical Learning: With Applications in Python. Springer International Publishing. <https://doi.org/10.1007/978-3-031-38747-0>

Juarez Pastor, L. (2022). Integration of the informal waste sector in the Indian city of Chennai: A case study. Application of the Institutional Network Analysis to a Municipal Waste Management System (Master's Thesis). TU Delft/Leiden University, MSc Industrial Ecology. <https://repository.tudelft.nl/islandora/object/uuid%3Ab1188b31-c14b-4c87-90ed-1b2c4ec8922f>

Kaufmann, B., Busby, D., Das, C. K., Tillu, N., Menon, M., Tewari, A. K., & Gorin, M. A. (2024). Validation of a Zero-shot Learning Natural Language Processing Tool to Facilitate

Data Abstraction for Urologic Research. *European Urology Focus*, S2405-4569(24)00012-9.
<https://doi.org/10.1016/j.euf.2024.01.009>

Khurana, D., Koli, A., Khatter, K., & Singh, S. (2023). Natural language processing: State of the art, current trends and challenges. *Multimedia Tools and Applications*, 82(3), 3713–3744.
<https://doi.org/10.1007/s11042-022-13428-4>

Klijn, F., De Bruijn, K. M., Knoop, J., & Kwadijk, J. (2012). Assessment of the Netherlands' Flood Risk Management Policy Under Global Change. *AMBIO*, 41(2), 180–192.
<https://doi.org/10.1007/s13280-011-0193-x>

Krefeld-Schwalb, A., Wei, S., & Gabel, S. (2024). Global Evidence on the Motives for Sustainable Behaviors. *OSF*. <https://doi.org/10.31234/osf.io/syku6>

KRO-NCRV. (2024). Honderden inwoners van Limburg eisen maatregelen tegen hoogwater.
<https://pointer.kro-ncrv.nl/honderden-inwoners-van-limburg-eisen-maatregelen-tegen-hoogwater>

Kuitert, G. (2021). Volt vooral populair onder studenten: 'Maar we zijn geen jongerenpartij'. *tubantia.nl*.
<https://www.tubantia.nl/enschede/volt-vooral-populair-onder-studenten-maar-we-zijn-geen-jongerenpartij~a9a9f52d/>

Lechner, J. (2022). Role of household climate change adaptation in reducing coastal flood risk: The case of Shanghai. (Master's thesis, Delft University of Technology).
<http://resolver.tudelft.nl/uuid:ccf10588-6a13-490a-a249-d866f5de3706>

Longford, N. T. (2008). Longitudinal and Time-Series Analysis. In: *Studying Human Populations*. Springer Texts in Statistics. Springer, New York, NY.
https://doi.org/10.1007/978-0-387-73251-0_11

Mesdaghi, B. (2020). Institutional interactions in climate adaptation of interdependent transport infrastructures: The case surrounding the Port of Rotterdam (Master's thesis). Delft University of Technology, Faculty of Technology, Policy and Management.
<https://repository.tudelft.nl/islandora/object/uuid%3A38eb2da8-c1cf-414f-a57c-753af89074f6>

Mesdaghi, B., Ghorbani, A., & de Bruijne, M. (2022). Institutional dependencies in climate adaptation of transport infrastructures: An Institutional Network Analysis approach. *Environmental Science & Policy*, 127, 120–136. <https://doi.org/10.1016/j.envsci.2021.10.010>

Metoui, N. & De Stefani, J. (2023) Responsible Data Analytics. Lectures. SEN163B. Complex Systems Engineering and Management. Technical University of Delft.

NL Times. (2021). Flood damage in Valkenburg estimated at €400 million; 700 families displaced. Retrieved from
<https://nltimes.nl/2021/07/21/flood-damage-valkenburg-estimated-eu400-million-700-families-displaced>

Noll, B., Filatova, T., & Need, A. (2022). Original Questionnaire for the 'SCALAR' household surveys on private climate change adaptation. Multi Actor Systems Department, Faculty of Technology, Policy and Management, Delft University of Technology, The Netherlands, and Faculty of Behavioral, Management and Social Sciences, University of Twente, The Netherlands. <https://doi.org/10.17026/dans-x9h-nj3w>

Noll, B., Filatova, T., Need, A., & Taberna, A. (2022). Contextualizing cross-national patterns in household climate change adaptation. *Nature Climate Change*, 12(1), 30–35. <https://doi.org/10.1038/s41558-021-01222-3>

Otjes, S. (2018). The phoenix of consensus democracy: Party system change in the Netherlands. In *Changing Party Systems in Western Europe*. Oxford University Press. <http://doi.org/10.4324/9781315147116-10>

Pavlova, M. K., & Lühr, M. (2023). Volunteering and political participation are differentially associated with eudaimonic and social well-being across age groups and European countries. *PLOS ONE*, 18(2), e0281354. <https://doi.org/10.1371/journal.pone.0281354>

Peppers, K., Tuunanen, T., Rothenberger, M. A., & Chatterjee, S. (2014). A Design Science Research Methodology for Information Systems Research. *MIS Quarterly*, 28(3), 45-77. <https://doi.org/10.2753/MIS0742-1222240302>

Rangapur, A., & Rangapur, A. (2024). The Battle of LLMs: A Comparative Study in Conversational QA Tasks. *arXiv*. <https://doi.org/10.48550/arXiv.2405.18344>

Robinson, S. B., & Leonard, K. F. (2018). *Designing Quality Survey Questions*. SAGE Publications. 978-1-5063-3053-2.

Runkler, T. A. (2016). *Data Analytics: Models and Algorithms for Intelligent Data Analysis*. ISBN: 9783658140748. <https://doi.org/10.1007/978-3-658-14075-5>

Schelling, T. C. (1971). Dynamic models of segregation. *The Journal of Mathematical Sociology*, 1(2), 143–186. <https://doi.org/10.1080/0022250X.1971.9989794>

Siddiki, S., Brady, U., & Frantz, C. K. (2023). The Institutional Grammar: Evolving Directions in Current Research. *International Review of Public Policy*, 5(2), Article 2. <https://doi.org/10.4000/irpp.3550>

De Silva, M. & Kawasaki, A. (2020). A local-scale analysis to understand differences in socioeconomic factors affecting economic loss due to floods among different communities. *International Journal of Disaster Risk Reduction*, 47, 101526. <https://doi.org/10.1016/j.ijdrr.2020.101526>

Snel, K. A. W., Hegger, D. L. T., Mees, H. L. P., Craig, R. K., Kammerbauer, M., Doorn, N., Bergsma, E., & Wamsler, C. (2022). Unpacking notions of residents' responsibility in flood risk governance. *Environmental Policy and Governance*, 32(3), 217-231. <https://doi.org/10.1002/eet.1985>

Törnberg, P. (2023). How to use LLMs for Text Analysis (arXiv:2307.13106). arXiv. <https://doi.org/10.48550/arXiv.2307.13106>

Verheul, D. (2021). "Candid climate efforts or empty promises?" Analyzing NSA's voluntary commitments and efforts (Master's thesis). Delft University of Technology, Faculty of Technology, Policy and Management. <https://repository.tudelft.nl/islandora/object/uuid%3A12a972fc-5358-4c7e-a477-1bf3fb3a5da0>

Wagenblast, T. (2022). Private Flood Adaptation and Social Networks: Using Agent-based Modelling to Explore the Effects of Private Flood Adaptation Policies in Presence of Social Networks and Information Diffusion. Delft University of Technology. <http://resolver.tudelft.nl/uuid:14e8a359-f0e9-40da-a838-09f2558b726e>

Watkins, C., & Westphal, L. M. (2016). People don't talk in institutional statements: A methodological case study of the institutional analysis and development framework. *Policy Studies Journal*, 44(S1), S98-S122. <https://doi.org/10.1111/psj.12139>

Wilke, C. (2019). *Fundamentals of Data Visualization: A Primer on Making Informative and Compelling Figures*. First edition. Sebastopol, CA: O'Reilly Media. <https://clauswilke.com/dataviz/>

Ziems, C., Held, W., Shaikh, O., Chen, J., Zhang, Z., & Yang, D. (2024). Can Large Language Models Transform Computational Social Science? *Computational Linguistics*. Advance online publication. <https://doi.org/10.48550/arXiv.2305.03514>

Appendix

A. Literature Review Problem Identification Overview

A1. Literature on Climate Change Adaptation

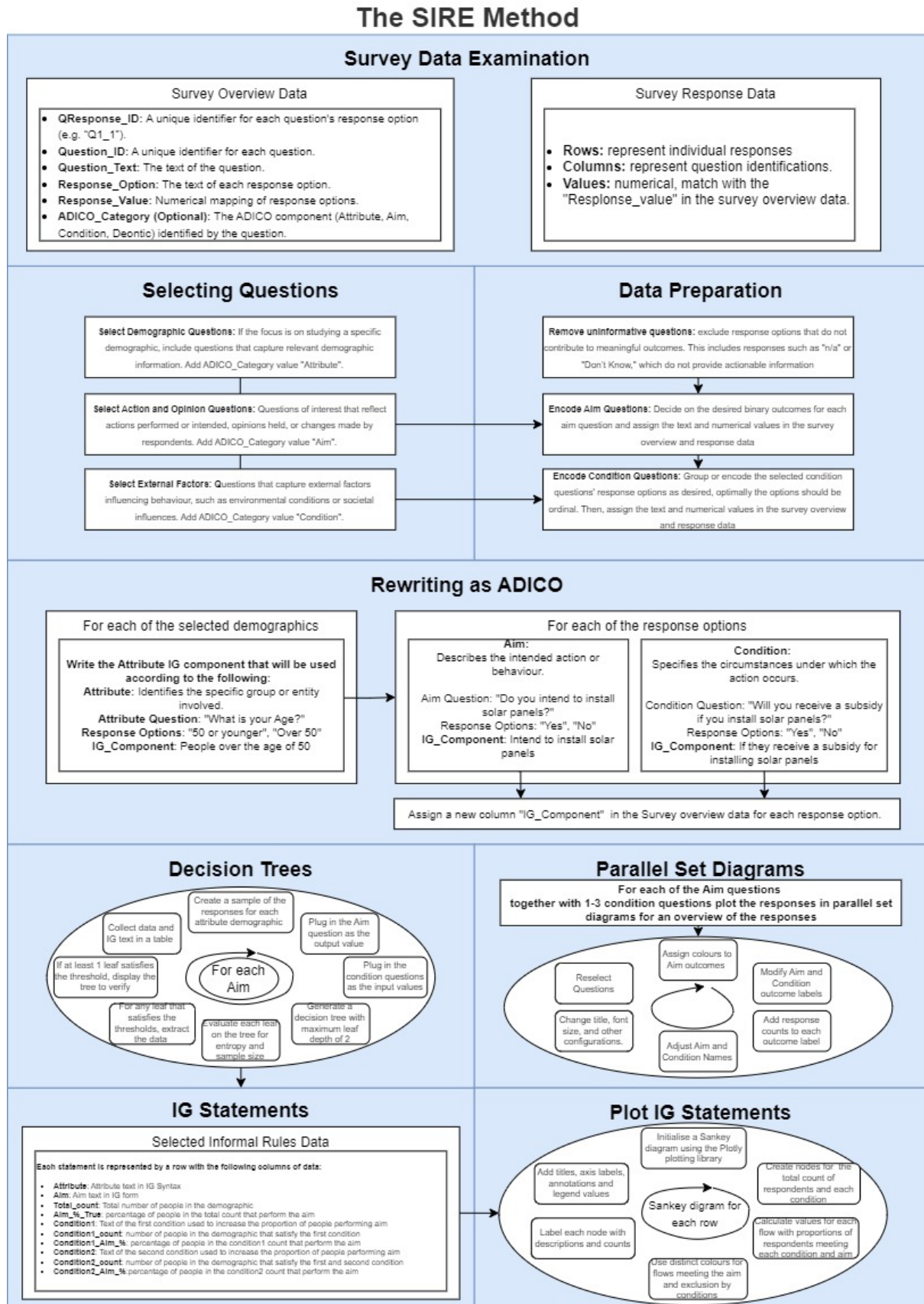
1. Wagenblast, T. (2022):
 - a. Topic: Private Flood Adaptation and Social Networks
 - b. Focus: Policies and norms within relevant regions, use of the dataset & agent-based modelling
 - c. Findings: Examines the interplay between private flood adaptation and public policies, emphasising the role of social networks. Highlights the significance of homogeneous networks in successful private adaptation and calls for future research on evolving social networks.
 - d. Limitations: Simplistic model representation, lack of consideration for policy impacts beyond information policies, and a need for broader case study testing.
2. Lechner, J. (2022):
 - a. Topic: Role of Household Climate Change Adaptation in Reducing Coastal Flood Risk
 - b. Focus: Policies and norms within relevant regions, using SCALAR dataset
 - c. Findings: Household climate change adaptation in Shanghai, emphasising the importance of quantifying household adaptation and addressing barriers like costs and regulations. Recommends policies addressing variations in adaptation behaviour among household groups.
 - d. Limitations: Data specificity, timing dependency, limited adaptation measures, behavioural assumptions, and the need for a more comprehensive model.
3. Abebe, Y. A. et al. (2020):
 - a. Topic: The Role of Household Adaptation Measures in Reducing Vulnerability to Flooding
 - b. Focus: Flood mitigation, agent-based modelling
 - c. Findings: Introduces a coupled agent-based and flood model to study individual adaptation behaviour and its influence on flood risk management. Emphasises the role of social networks and simple measures in reducing vulnerability.
 - d. Limitations: Simplified decision trees, model configuration concerns, need for intelligent decision-making models, flood model refinement, and generalizability challenges.
4. De Silva, M. et al. (2020):
 - a. Topic: Socio-economic Impacts of Floods
 - b. Focus: Analyses the socioeconomic influences on economic loss due to floods in Rathnapura, Sri Lanka.
 - c. Findings: Reveals the direct impact of flood characteristics and household income on economic loss, emphasising the need to consider diverse economic groups in disaster risk management.
 - d. Limitations: No specific limitations are mentioned in the provided text.

5. Chan, F. K. S. et al. (2022):
 - a. Topic: Comparison of Sustainable Flood Risk Management in Various Countries
 - b. Focus: Policymaking in various countries, flood risk mitigation
 - c. Findings: Reviews sustainable flood risk management practices in the UK, the Netherlands, the US, and Japan. Emphasises the importance of integrating sustainability concepts and moving away from traditional hard-engineering approaches.
 - d. Limitations: Acknowledges different interpretations of SFRM and recommends a context-specific approach.
6. Erdlenbruch, K. et al. (2018):
 - a. Topic: Flood Mitigation, Policy & Simulation
 - b. Focus: Simulation of individual adaptation dynamics, communication policies
 - c. Findings: Provides insights into the diffusion of individual adaptation measures, evaluates communication policies, and identifies key dynamic parameters. Recommends more studies on barriers to implementation and sustained use of adaptation measures.
 - d. Limitations: Discusses potential extensions to the model, including mobile households, flood event dependence, different institutional settings, and economic incentives.
7. Haer, T. et al. (2016):
 - a. Topic: Flood-Risk Mitigation Analysis Using Agent-Based Modeling
 - b. Focus: Effectiveness of flood risk communication strategies
 - c. Findings: The effectiveness of flood risk communication strategies using an agent-based modelling approach. Highlights the importance of tailored communication and leveraging social networks for more effective risk communication.
 - d. Limitations: Emphasises the need for tailoring communication strategies to the local context and obtaining empirical data for modelling specific regions.
8. Snel, K. A. W. et al. (2022):
 - a. Topic: Individual Responsibility, Socio-economic Role & Impacts, Flood Risk Governance
 - b. Focus: Examines theoretical notions of individual responsibility in flood-prone areas and analyses flood risk governance practices across countries.
 - c. Findings: Explores connections and differences between individuals' perceptions of responsibility and government assumptions. Provides insights into the complex nature of individual responsibility in flood risk governance.
 - d. Limitations: Acknowledges cultural differences in perception, calls for in-depth country-specific analyses, and recommends future research on responsibility divisions in other institutional settings.

A2. Literature Using the Institutional Network Analysis

1. Mesdaghi, B. (2020):
 - a. Topic: Institutional interactions in climate adaptation of interdependent transport infrastructures
 - b. Focus: Examines institutional interactions within the transport infrastructures connected to the Port of Rotterdam, employing methodologies such as desk research and semi-structured interviews.
 - c. Findings: The study highlights the importance of distinguishing between actual and ideal institutional practices and suggests using agent-based modelling for a better understanding of dynamic institutional interactions.
 - d. Limitations: Points to a lack of depth in exploring the impacts of drought and heat, limited understanding of infrastructure interdependencies, and the constraints of online interviews limiting observational data. Also notes bias in data coding and clustering.
2. Verheul, D. (2021):
 - a. Topic: Non-state actors' voluntary commitments to climate change mitigation
 - b. Focus: Analyses voluntary commitments of food and beverage companies to climate initiatives like the Science-Based Targets, using Institutional Network Analysis to assess participation dynamics.
 - c. Findings: Identifies socio-economic constructs and societal pressure as key mechanisms influencing corporate behaviour towards climate change mitigation. Recommends enhancing data collection and analysis techniques to better understand these dynamics.
 - d. Limitations: Discusses limitations such as the choice of institutional over neoclassical economics, the challenge of non-public data, online interviews limiting non-verbal cues, and the need for broader theoretical discussions with interviewees to enhance understanding.
3. Juarez Pastor, L. (2022)
 - a. Topic: Integration of the informal waste sector in Chennai's waste management system
 - b. Focus: Investigates the role of the informal waste sector within Chennai's municipal waste management using Institutional Network Analysis.
 - c. Findings: Identifies major bottlenecks such as the municipality's reluctance to integrate informal waste workers and highlights the gaps between formal rules and actual practices. Suggests that a clearer policy framework and better stakeholder engagement could facilitate integration.
 - d. Limitations: Points out issues like the underrepresentation of certain stakeholders in data collection, cultural and contextual disconnections affecting analysis, and ambiguities in institutional grammar that might compromise data integrity. Recommends a more dynamic and comparative approach to understand and resolve these issues.

B. SIRE Method Diagram



C. Standardised Python Scripts for the SIRE Method

Package installations for Python:

```
• pandas==2.1.4
• numpy==1.26.3
• langchain_core==0.1.52
• langchain_groq==0.1.3
• plotly==5.20.0
• sklearn==1.3.2
• matplotlib==3.8.2
```

C1. Survey Data Examination

```
#Change these locations to your Overview and Responses csv files:
Survey_Overview_Location = "..\Generate mock survey\Mock_Survey_Overview.csv"
Survey_Responses_Location = "..\Generate mock survey\mock_survey_responses.csv"

#Reads the files into a table
Survey_Overview = pd.read_csv(Survey_Overview_Location)
Survey_Responses = pd.read_csv(Survey_Responses_Location).set_index('Response_ID')

#creates a unique id value for each response option
Survey_Overview["Question_Response_ID"] = Survey_Overview["Question_ID"].astype(str) + "_"
+ Survey_Overview["Response_Value"].astype(str)
Survey_Overview.set_index('Question_Response_ID', inplace=True)
```

```
# You can run this if you have run everything previously and have saved a completed
overview file
Completed_Overview_Location = "Mocksurvey_completed_overview.csv" #change the name to the
name of the file

try:
    #Reads the files into a table
    adjusted_overview =
pd.read_csv(Completed_Overview_Location).set_index('Question_Response_ID')
    #This way this file can be loaded and various steps of adding columns to the overview
can be skipped
    display(adjusted_overview)
except:
    print("The requested completed overview file does not exist or can not be found. If you
know you have saved it before, please verify the name and location")
```

C2. Selecting and Categorising Questions

```
try: #Make lists of questions that have been categorised into each ADICO component:
    attributes = Survey_Overview[Survey_Overview['ADICO_Category'] ==
"Attribute"]['Question_ID'].unique()
    aims = Survey_Overview[Survey_Overview['ADICO_Category'] ==
"Aim"]['Question_ID'].unique()
    conditions = Survey_Overview[Survey_Overview['ADICO_Category'] ==
"Condition"]['Question_ID'].unique()
```

```

deontics = Survey_Overview[Survey_Overview['ADICO_Category'] ==
"Deontic"]['Question_ID'].unique()
except Exception as e: print(e) #This won't work if you don't have the 'ADICO_Category'
column

#Alternatively you can make manual lists of the questions if you wish to select a subset or
if the questions have not been categorised with ADICO_Category:
aims = ['Q3', 'Q4', 'Q5']
conditions = ['Q1', 'Q2', 'Q6', 'Q7', 'Q8'] #Added attribute questions to conditions
deontics = ['Q9', 'Q10']

```

C3. Data Preparation and Preprocessing

```

adjusted_responses = Survey_Responses.copy().dropna(axis=1)

#Prepare responses to be binary outcomes
adjusted_responses.loc[adjusted_responses["Q1"] >= 4, "Q1"] = 4 #sets all responses of age
over 50 to 4
adjusted_responses.loc[adjusted_responses["Q1"] <= 3, "Q1"] = 3 #sets all responses of age
under 50 to 3

adjusted_responses.loc[adjusted_responses["Q7"] >= 4, "Q7"] = 4 #sets all responses that
consider public transport price expensive to 4
adjusted_responses.loc[adjusted_responses["Q7"] <= 3, "Q7"] = 3 #sets all responses that do
not consider public transport price expensive to 3

#Remove responses that might not be relevant:
#Note: this is just an example on how to filter for a subset of the responses, ignoring
"other" as a gender may be ethically problematic and misrepresent true patterns of
behaviour
adjusted_responses = adjusted_responses[adjusted_responses["Q2"] != 3] #This removes all
"other" responses to the question on gender

```

```

#You should skip this step if you have loaded an adjusted and completed overview file

adjusted_overview = Survey_Overview.copy().dropna(axis=1)

#Prepare responses to be binary outcomes
adjusted_overview.loc[(adjusted_overview["Question_ID"] == "Q1") &
(adjusted_overview["Response_Value"] == 4), "Response_Option"] = "Over 50" #changes the
text of the response for age option 4 to "Over 50"
adjusted_overview.loc[(adjusted_overview["Question_ID"] == "Q1") &
(adjusted_overview["Response_Value"] == 3), "Response_Option"] = "50 or younger" #changes
the text of the response for age option 3 to "50 or younger"

adjusted_overview.loc[(adjusted_overview["Question_ID"] == "Q7") &
(adjusted_overview["Response_Value"] == 4), "Response_Option"] = "Expensive" #changes the
text of the response for price perception option 4 to "Expensive"
adjusted_overview.loc[(adjusted_overview["Question_ID"] == "Q7") &
(adjusted_overview["Response_Value"] == 3), "Response_Option"] = "not Expensive" #changes
the text of the response for price perception option 3 to "not Expensive"

#If it doesn't exist yet, create a column in overview which should eventually contain the
question and response rewritten in ADICO syntax
try: adjusted_overview["IG_Component"]

```

```
except: adjusted_overview["IG_Component"] = adjusted_overview["Question_Text"] + " Response:"
" + adjusted_overview["Response_Option"]
```

C4. Rewriting to ADICO components (LLM API Example)

```
#Ask Generative AI to rewrite the Jsn Approach
#groq offers these models:
models = ["gemma-7b-it", "llama3-8b-8192", "mixtral-8x7b-32768", "llama3-70b-8192"]
# Define the LLM, models[1] represents "llama3-8b-8192"
llm = ChatGroq(temperature=0, model=models[1], api_key=groqkey)

#if you prefer to use openAI (requires paid API access):
# models = ["gpt-3.5-turbo-0125", "gpt-4o"]
# llm = ChatOpenAI(temperature=0, model=models[1], api_key=OPENAI_Key)

#You can experiment with the system instructions given to the ai, this section was excluded
because I considered it redundant but it can always be added
"""Here is an example: "Attribute": "People", "Aim": "do this specific action",
"Condition1": "if this condition is met", "Condition2": "and this condition is met"
I should be able to combine any output Attribute + Aim + Condition to form a full
3rd-person sentence that describes behaviour. """

#Function that gets the response from the LLM model and reads it into a table, the function
takes a dataframe table of the survey overview,
# preferably provide it with a subset of the survey overview of Aims, Conditions or
Attributes and indicate this in the ADICO_component value.
def ExampleCompletionFunction(questions, ADICO_component):
    #create an empty table with the output column:
    IG_components = pd.DataFrame(columns=["IG_Component"])

    #Handle each question separately to avoid confusion:
    for question in questions['Question_ID'].unique():

        # Prepares the provided overview table for the LLM
        questions = questions[questions['Question_ID'] ==
question][["Question_Text", "Response_Option"]]
        request = questionset.to_json(orient='index', index=True)

        #The general instructions for the LLM
        system = f"""
You are a json interpreter that transforms survey questions and responses into structured
informal rule institutional statement components in 3rd person.
The input json will have the following structure:
Object Key = identifier of question and answer pair
With nested values
Question_Text: Text of the question being asked
Response_Option = a response option for the question

You output json must have the following structure:
Object Key = the same identifier of question and answer pair
IG_Component = The rewritten question and answer

Be concise but do not simplify or generalise the actions and conditions.
If a response option is just a number, relate it to the other options for that
questionAlways produce a single JSON containing the same (number of) question identifiers
as have been provided.
```

In this instance the questions you have been provided are '{ADICO_component}'. Perform the text conversion below."""

#Additional information for the LLM that depends on the "ADICO_component" given:

```
system += {
'Attributes': 'demographic of the survey responses, should be written as "People (rewritten question and answer)" e.g. "People who are male"', #if "ADICO_component" is
'Attributes': 'Aims': "action question and the response, should be written as an action that precedes the subject of the sentence e.g. 'sell their house', do not include the subject of the sentence", #if "ADICO_component" is Aims
'Conditions': "condition question and response, should be rewritten as 'if (rewritten question and answer)' in 3rd person ('they') e.g. 'if they feel they cannot rely on government support'"}[ADICO_component] #if "ADICO_component" is Conditions
```

#Combines and formats the message and system instructions for the LLM

human = "{text}"

prompt = ChatPromptTemplate.from_messages([("system", system), ("human", human)])

#Sends the request to the LLM to get a response

chain = prompt | llm

response = chain.invoke({"text": request})

Extract the JSON string from the LLM response text

start = response.content.find("{")

end = response.content.rfind("}") + 1

json_str = response.content[start:end]

Parse the JSON string into a Python dictionary

try:

data_dict = json.loads(json_str)

except:

Replace all instances of "}}" with "}", and remove any incorrect json formatting

json_str = json_str.replace("}}", "}")+"}"

json_str = json_str.replace(",\n(", ",\n")

json_str = json_str.replace(",{", "{,")

Load the JSON string into a dictionary

data_dict = json.loads(json_str)

Read the processed response as a dataframe

IG_components = pd.concat([IG_components, pd.DataFrame.from_dict(data_dict, orient='index')])

return IG_components

Update adjusted_overview with values from the LLM's response

adjusted_overview.update(ExampleCompletionFunction(adjusted_overview[adjusted_overview['Question_ID'].isin(aims)], "Aims"))

adjusted_overview.update(ExampleCompletionFunction(adjusted_overview[adjusted_overview['Question_ID'].isin(attributes)], "Attributes"))

adjusted_overview.update(ExampleCompletionFunction(adjusted_overview[adjusted_overview['Question_ID'].isin(conditions)], "Conditions"))

```
# for now we treat deontic example questions as Conditions as deontics have not yet been
integrated in the final method
adjusted_overview.update(ExampleCompletionFunction(adjusted_overview[adjusted_overview['Que
stion_ID'].isin(deontics)], "Conditions"))
```

```
#View the updated table and see IG_component:
display(adjusted_overview)
#note that Attributes have been treated as conditions because they were included in the
list of conditions

#Save the table with rewritten IG components to a csv:
adjusted_overview.to_csv("Mocksurvey_completed_overview.csv", index=True)
#This way the overview file will all added columns can be loaded and previous steps can be
skipped
```

C5. Create Parallel Set Diagrams

```
# Select which aim and condition questions you would like to visualise in the Parallel set
diagram
aim = "Q4"
conditions = ["Q1", "Q7", "Q10"]
questions = conditions + [aim]

parallelsetdf = adjusted_responses[questions].astype(str)

# Define custom color mapping
color_mapping = {
    "1": "#00a6d6",
    "2": "#a7a7a7"
}

# Apply color mapping
parallelsetdf['color'] = parallelsetdf[aim].map(color_mapping)
total_responses = len(parallelsetdf)
aim_percentages = (parallelsetdf[aim].value_counts() / total_responses * 100).round(1)

# Update labels with counts and percentage
aimcounts = parallelsetdf[aim].value_counts()
parallelsetdf.loc[parallelsetdf[aim] == str(1), aim] = f"Do<br>n =
{aimcounts.loc[str(1)]}<br>({aim_percentages.loc[str(1)]}%"
parallelsetdf.loc[parallelsetdf[aim] == str(2), aim] = f"Do not<br>n =
{aimcounts.loc[str(2)]}<br>({aim_percentages.loc[str(2)]}%"

for condition in conditions:
    conditioncounts = parallelsetdf[condition].value_counts()
    for val in parallelsetdf[condition].unique():
        # Calculate the percentage of people that do the aim within the condition value
        condition_percentage =
(len(parallelsetdf[(parallelsetdf[condition]==val)&(parallelsetdf["color"]=="#00a6d6")]) /
conditioncounts.loc[val] * 100).round(1)
        # Update condition labels
        label = adjusted_overview.loc[(adjusted_overview['Question_ID'] == condition) &
(adjusted_overview['Response_Value'] == float(val)), "Response_Option"].iloc[0]
        parallelsetdf.loc[parallelsetdf[condition] == val, condition] = f"{label}<br>n =
{conditioncounts.loc[val]}<br>Do: {condition_percentage}%"
```

```

diagram_headers = ["Age?", "Perceived Price", "Peer pressure to use?", "Frequent use?"]
parallelsetdf.columns = diagram_headers + ["color"]

# Create the parallel categories plot
fig = px.parallel_categories(
    parallelsetdf,
    dimensions=diagram_headers,
    color='color',
)

# Update the layout of the figure with a title and font size
fig.update_layout(title_text="People: " +
adjusted_overview[adjusted_overview['Question_ID'] == aim]["Question_Text"].iloc[0],
                    height=500,
                    width=1050,
                    title_font=dict(size=30), # Set title font size here
                    font=dict(size=17) # Set default font size here (for labels, etc.)
)

# Show the plot
fig.show()

```

C6. Extracting Statements Using Decision Trees

```

# Function that takes a tree and an index value of a tree node that satisfies the
conditions to be a statement and generates a row for the statement table
def makeTreeStatement(tree, satisfactory_index, classes, features, attributeGroup, aim,
Aim_description):
    # Get the aim outcome of the node (which action has more samples)
    class_index = np.argmax(tree.value[satisfactory_index])
    Aim_resp = classes[class_index]

    # Retrieve the class counts at the root node
    aim_resp_count = tree.value[0][0] # Gets the counts for each class

    # get how many in people answered with this aim outcome
    total_aim_percent = aim_resp_count[class_index]/tree.n_node_samples[0]

    if satisfactory_index == 0:
        condition2 = condition1 = None,
        Condition2_description = Condition1_description = None
        condition2_rows_of_leaf = condition1_rows_of_leaf = []
        condition2_count = condition1_count = tree.n_node_samples[0]
        condition2_aim_resp_percent = condition1_aim_resp_percent = total_aim_percent
    else:
        condition1 = features[tree.feature[0]] # Get the condition used for the root node
        #get the description of the first condition question used
        Condition1_description = adjusted_overview[adjusted_overview["Question_ID"] ==
condition1.strip()]['Question_Text'].iat[0]
        #Rows of question data for condition1
        condition1_rows = adjusted_overview[(adjusted_overview['Question_ID'] ==
condition1) &
adjusted_overview['Response_Value'].isin(adjusted_responses[condition1].unique())]

        satisfactory_leaf_is_left = tree.feature[satisfactory_index - 1] != -2 # Determine
if the leaf is on the "left"

```

```

        if satisfactory_leaf_is_left: # The leaf meets the threshold
            leaf_parent_index = satisfactory_index - 1
        else: # The leaf does not meet the threshold (falls outside the threshold)
            leaf_parent_index = satisfactory_index - 2
            if tree.feature[leaf_parent_index] == -2: #was directly connected to root and
on the right
                leaf_parent_index = 0

        if leaf_parent_index != 0: #if the leaf is not directly connected to the root node
            root_threshold = tree.threshold[0] # get the threshold used on the root node
            parent_is_left = leaf_parent_index - 1 == 0 # Determine if the parent is on
the "left"

            # Set 'in_thresh' based on whether the condition is met (leaf is within the
threshold)
            if parent_is_left: # The parent meets the threshold
                condition1_rows_of_leaf = condition1_rows[condition1_rows['Response_Value']
<= root_threshold]['IG_Component']

            else: # The leaf does not meet the threshold
                condition1_rows_of_leaf = condition1_rows[condition1_rows['Response_Value']
> root_threshold]['IG_Component']

            condition1_count = tree.n_node_samples[leaf_parent_index]
            condition1_aim_resp_percent = tree.value[leaf_parent_index][0,
class_index]/condition1_count
            condition2 = features[tree.feature[leaf_parent_index]]
            Condition2_description = Survey_Overview[Survey_Overview["Question_ID"] ==
condition2.strip()]['Question_Text'].iat[0]

            threshold2 = tree.threshold[leaf_parent_index] # get the threshold used on the
least entropy leaf
            #Rows of question data for this condition
            condition2_rows = adjusted_overview[(adjusted_overview['Question_ID'] ==
condition2) &
adjusted_overview['Response_Value'].isin(adjusted_responses[condition2].unique())]

            if satisfactory_leaf_is_left:
                condition2_rows_of_leaf = condition2_rows[condition2_rows['Response_Value']
<= threshold2]['IG_Component']

            else:
                condition2_rows_of_leaf = condition2_rows[condition2_rows['Response_Value']
> threshold2]['IG_Component']

            condition2_count = tree.n_node_samples[satisfactory_index]
            condition2_aim_resp_percent =
tree.value[satisfactory_index][0][np.argmax(tree.value[satisfactory_index][0])]/condition2_
count

        else:
            condition1 = features[tree.feature[leaf_parent_index]]
            parent_threshold = tree.threshold[leaf_parent_index] # get the threshold used
on the least entropy leaf
            #Rows of question data for this condition

```

```

        condition1_rows = adjusted_overview[(adjusted_overview['Question_ID'] ==
condition1) &
adjusted_overview['Response_Value'].isin(adjusted_responses[condition1].unique())]

        if satisfactory_leaf_is_left:
            condition1_rows_of_leaf = condition1_rows[condition1_rows['Response_Value']
<= parent_threshold]['IG_Component']
        else:
            condition1_rows_of_leaf = condition1_rows[condition1_rows['Response_Value']
> parent_threshold]['IG_Component']

        condition1_count = tree.n_node_samples[satisfactory_index]
        condition1_aim_resp_percent =
tree.value[satisfactory_index][0][np.argmax(tree.value[satisfactory_index][0])]/condition1_
count

        condition2 = None,
        Condition2_description = None
        condition2_rows_of_leaf = []
        condition2_count = condition1_count
        condition2_aim_resp_percent = condition1_aim_resp_percent

    # Create a new row with specified values
    new_row = {
        'Attribute': attributeGroup,

        'Aim': aim,
        'Aim_description': Aim_description,
        'Aim_resp': Aim_resp,    # Class with the lowest entropy
        'Total_count': tree.n_node_samples[0],
        'Aim_%_True': total_aim_percent,

        'Condition1': condition1,
        'Condition1_description': Condition1_description,
        'Condition1_resp': ", ".join(condition1_rows_of_leaf),
        'Condition1_count': condition1_count,
        'Condition1_Aim_%': condition1_aim_resp_percent,

        'Condition2': condition2,
        'Condition2_description': Condition2_description,
        'Condition2_resp': ", ".join(condition2_rows_of_leaf),
        'Condition2_count': condition2_count,
        'Condition2_Aim_%': condition2_aim_resp_percent,

        'final_entropy': tree.impurity[satisfactory_index]
    }
    return new_row

```

```

def select_statements(aim, conditions, responses, statement_questions, attributeGroup,
entropy_threshold, sample_threshold, recursion):

    # Find the relevant row and value labels
    aimRow = adjusted_overview[adjusted_overview["Question_ID"] == aim]
    aimRow0 = str(aimRow.at[aimRow.index[0], "IG_Component"])
    aimRow1 = str(aimRow.at[aimRow.index[-1], "IG_Component"])
    classes = [aimRow0, aimRow1]

```

```

    Aim_description = adjusted_overview[adjusted_overview["Question_ID"] ==
aim.strip()]['Question_Text'].iat[0]

    print(f"Processing aim: {aim}: {str(aimRow.at[aimRow.index[0], 'Question_Text'])}")

    # Split the data
    condition_responses = responses[conditions].values
    aim_responses = responses.loc[:, aim].values

    sample_demographic_size = len(aim_responses)

    clf_entropy = DecisionTreeClassifier(criterion='entropy', random_state=42, max_depth=2)
    clf_entropy.fit(condition_responses, aim_responses)

    #list of the conditions we test
    features = conditions

    # 'clf_entropy' is your trained decision tree classifier
    tree = clf_entropy.tree_

    #Identify leaf indices (where feature is -2)
    leaf_indices = np.where(tree.feature == -2)[0]

    #Find nodes where entropy is within the threshold
    pure_enough_nodes = np.where(tree.impurity <= entropy_threshold)[0]
    #Find nodes which have enough samples is within the threshold
    enough_sample_nodes = np.where(tree.n_node_samples >=
int(sample_threshold*sample_demographic_size))[0]
    # Find the intersection of both lists
    satisfactory_nodes = np.intersect1d(pure_enough_nodes, enough_sample_nodes)

    for node in satisfactory_nodes:
        new_row = makeTreeStatement(tree, node, classes, features, attributeGroup, aim,
Aim_description)
        # Using a new index to add a row directly
        row_number = len(statement_questions) # Determine the next index
        statement_questions.loc[row_number] = new_row
        # Drop duplicates if needed
        statement_questions = statement_questions.drop_duplicates()

    #find the leaf with the lowest entropy
    least_entropy_index = leaf_indices[np.argmin(tree.impurity[leaf_indices])]
    final_entropy = tree.impurity[least_entropy_index]
    print(final_entropy)

    if final_entropy < entropy_threshold:
        # Visualize the decision tree
        plt.figure(figsize=(25, 5))

        plot_tree(clf_entropy, filled=True, feature_names=features, class_names=classes)
        plt.suptitle(f"{aim}: {str(aimRow.at[aimRow.index[0], 'Question_Text'])}")
        plt.show()

    if recursion < min(3, len(conditions)-2):
        recursion_conditions = [condition for condition in conditions if condition not in
statement_questions["Condition1"].tolist()[~recursion:]]

```

```

        statement_questions = select_statements(aim, recursion_conditions, responses,
statement_questions, attributeGroup, entropy_threshold, sample_threshold, recursion+1)

    return statement_questions

```

```

#create a table that will be filled by selected combinations of aims and conditions with
the relevant information for further processing
statement_questions = pd.DataFrame(columns = ['Attribute',
                                             'Aim', 'Aim_description', 'Aim_resp',
                                             'Total_count', 'Aim_%_True',
                                             'Condition1', 'Condition1_description',
                                             'Condition1_resp', 'Condition1_count', 'Condition1_Aim_%',
                                             'Condition2', 'Condition2_description',
                                             'Condition2_resp', 'Condition2_count', 'Condition2_Aim_%',
                                             'final_entropy'])

```

```

#You can change these values:
entropy_threshold = 0.85 #How much entropy can the node have (yes aim:no aim) to be
considered?
sample_threshold = 0.1 #What proportion of the original sample size to be considered?

# Select which aim and condition questions you would like to visualize in the Parallel set
diagram
aim = "Q4"
conditions = ["Q1", "Q7", "Q10"]

statement_questions = select_statements(aim, conditions, adjusted_responses,
statement_questions, "People", entropy_threshold, sample_threshold, 0)

display(statement_questions)

# Save the DataFrame to a CSV file
statement_questions.to_csv("Mocksurvey_tree_selected_statements.csv", index=False)

```

C7. Visualising Extracted Statements

```

import plotly.graph_objects as go

def insert_linebreak(text):
    words = text.split()
    if len(words) > 8:
        middle_index = len(words) // 2
        words[middle_index] = words[middle_index] + "<br>"
        return " ".join(words)
    return text

for i, row in statement_questions.iterrows():
    # Calculate values
    total_count = row['Total_count']
    condition1_count = row['Condition1_count']
    condition2_count = row['Condition2_count']

    aim_true_total = total_count * row['Aim_%_True']
    aim_false_total = total_count - aim_true_total

    aim_true_cond1 = condition1_count * row['Condition1_Aim_%']

```

```

aim_false_cond1 = condition1_count - aim_true_cond1
excluded_true_cond1 = aim_true_total - aim_true_cond1
excluded_false_cond1 = aim_false_total - aim_false_cond1

aim_true_cond2 = condition2_count * row['Condition2_Aim_%']
aim_false_cond2 = condition2_count - aim_true_cond2
excluded_true_cond2 = aim_true_cond1 - aim_true_cond2
excluded_false_cond2 = aim_false_cond1 - aim_false_cond2

# Define node labels with counts of leaving flows
labels = [
    f"Applies: {round(aim_true_total)}{(' <br>' * round(max(aim_true_total, aim_false_total) / total_count * 24 + 1))}Does not apply: {round(aim_false_total)}",
    f"{round(aim_true_cond1)}{(' <br>' * round((aim_true_cond1 + excluded_true_cond1) / total_count * 12 + 1))}{round(excluded_true_cond1)}",
    f"{round(aim_false_cond1)}{(' <br>' * round((aim_false_cond1 + excluded_false_cond1) / total_count * 12 + 1))}{round(excluded_false_cond1)}",
    f"{round(aim_true_cond2)}{(' <br>' * round((aim_true_cond2 + excluded_true_cond2) / total_count * 10 + 1))}{round(excluded_true_cond2)}",
    f"{round(aim_false_cond2)}{(' <br>' * round((aim_false_cond2 + excluded_false_cond2) / total_count * 10 + 1))}{round(excluded_false_cond2)}",
    f"Applies: {aim_true_cond2:.0f}",
    f"Does not apply: {aim_false_cond2:.0f}",
    f"Excluded: {excluded_true_cond1 + excluded_false_cond1 + excluded_true_cond2 + excluded_false_cond2:.0f}"
]

# Define links
source = [0, 0, 1, 2, 1, 2, 3, 4, 3, 4]
target = [1, 2, 3, 4, 7, 7, 5, 6, 7, 7]
value = [
    aim_true_total, aim_false_total,
    aim_true_cond1, aim_false_cond1, excluded_true_cond1, excluded_false_cond1,
    aim_true_cond2, aim_false_cond2, excluded_true_cond2, excluded_false_cond2,
]

# Colour scheme
color_aim_true = '#00a6d6'
color_aim_false = '#a7a7a7'
color_excluded = 'grey'

link_colors = [color_aim_true, color_aim_false,
               color_aim_true, color_aim_false, color_excluded, color_excluded,
               color_aim_true, color_aim_false, color_excluded, color_excluded]

# Create Sankey diagram
fig = go.Figure(data=[go.Sankey(
    arrangement = "perpendicular",
    node = dict(
        pad = 50,
        thickness = 20,
        line = dict(width = 0.5),

```

```

        label = labels,
        colour = "blue",
    ),
    link = dict(
        source = source,
        target = target,
        value = value,
        colour = link_colors
    )))

# Add invisible scatter traces for legend
fig.add_trace(go.Scatter(
    x=[None], y=[None], mode='markers',
    marker=dict(size=15, colour='#00a6d6'),
    name='Aim applies',
    showlegend=True
))

fig.add_trace(go.Scatter(
    x=[None], y=[None], mode='markers',
    marker=dict(size=15, colour='#a7a7a7'),
    name='Aim does not apply',
    showlegend=True
))

fig.add_trace(go.Scatter(
    x=[None], y=[None], mode='markers',
    marker=dict(size=15, colour='grey'),
    name='Excluded by conditions',
    showlegend=True,
))

# Add annotations below the diagram, format/wrap the text so that it fits
annotations = [
    dict(x=0, y=-0.05, align="center", xref="paper", yref="paper",
text=row['Attribute'], showarrow=False),
    dict(x=0.33, y=-0.075, align="center", xanchor="center", xref="paper",
yref="paper", text=insert_linebreak(row['Condition1_resp']), showarrow=False),
    dict(x=0.66, y=-0.075, align="center", xanchor="center", xref="paper",
yref="paper", text=insert_linebreak(row['Condition2_resp']), showarrow=False),
    dict(x=1, y=-0.05, align="center", xref="paper", yref="paper", text="final",
showarrow=False)
]

# Update layout
fig.update_layout(
    title_text=f""""Sankey plot representation of an IG statement extracted by decision
tree algorithm:<br>{row['Aim_resp']}""",
    font_size=20,
    height=700, # Increased height to accommodate annotations
    width=1500,
    legend=dict(
        orientation="h",
        yanchor="bottom",
        y=1,
        xanchor="right",
        x=1

```

```

),
# Hide axes
xaxis=dict(visible=False),
yaxis=dict(visible=False),
# Remove plot background
plot_bgcolor='rgba(0,0,0,0)',
# Add annotations
annotations=annotations
)

fig.show()

```

D. Example Outputs of the SIRE Method Python Script

In this Appendix example data, visualisations and outputs are presented to help readers understand what data format the SIRE method requires and what information it can produce. All examples are based on a mock survey that was inspired from the SCALAR and ESS datasets. The mock survey data does not include real responses and is not representative of real opinions. The data is merely standardised to be optimised for the SIRE method.

D1. Survey Data Examination and Categorising Questions

Table D.1: Example Survey Overview Data

Question_ID	Question_Text	Response_Option	Response_Value	ADICO_Category
Q1	What is your age?	Under 18	1	Attribute
Q1	What is your age?	18-35	2	Attribute
Q1	What is your age?	36-50	3	Attribute
Q1	What is your age?	51-65	4	Attribute
Q1	What is your age?	Over 65	5	Attribute
Q2	What is your gender?	Male	1	Attribute
Q2	What is your gender?	Female	2	Attribute
Q2	What is your gender?	Other	3	Attribute
Q3	Have you recycled in the past month?	Yes	1	Aim
Q3	Have you recycled in the past month?	No	2	Aim
Q4	Do you use public transportation frequently?	Yes	1	Aim
Q4	Do you use public transportation frequently?	No	2	Aim
Q5	Did you vote in the last election?	Yes	1	Aim
Q5	Did you vote in the last election?	No	2	Aim
Q6	Do you think recycling is easy?	Yes	1	Condition
Q6	Do you think recycling is easy?	No	2	Condition

Q7	Do you consider public transport to be affordable?	Very affordable	1	Condition
Q7	Do you consider public transport to be affordable?	Affordable	2	Condition
Q7	Do you consider public transport to be affordable?	Neutral	3	Condition
Q7	Do you consider public transport to be affordable?	Expensive	4	Condition
Q7	Do you consider public transport to be affordable?	Very expensive	5	Condition
Q8	Do you follow political news on Social media?	Yes	1	Condition
Q8	Do you follow political news on Social media?	No	2	Condition
Q9	Do you feel obligated to recycle?	Yes	1	Deontic
Q9	Do you feel obligated to recycle?	No	2	Deontic
Q10	Do your peers expect you to use public transportation?	Yes	1	Deontic
Q10	Do your peers expect you to use public transportation?	No	2	Deontic

Table D.2: Example Survey Responses Data

Response_ID	Q1	Q2	Q3	Q4	Q5	Q6	Q7	Q8	Q9	Q10
1	5	3	2	2	2	1	2	1	2	1
2	3	1	1	1	1	1	4	2	1	1
3	1	3	1	2	1	1	1	2	1	1
4	1	1	2	2	1	2	4	1	1	2
5	4	1	2	2	2	1	1	2	2	1
6	2	1	1	1	1	1	2	1	2	2
7	3	3	2	1	1	2	1	1	1	2
8	5	1	2	1	2	1	3	1	1	2
9	1	2	2	1	2	2	3	1	1	2
10	5	2	1	2	1	1	1	1	2	1

D2. Data Preparation and Preprocessing

Table D.3: Processed Survey Overview Values

Question_ID	Question_Text	Response_Option	Response_Value	ADICO_Category
Q1	What is your age?	50 or younger	3	Attribute
Q1	What is your age?	Over 50	4	Attribute
Q2	What is your gender?	Male	1	Attribute
Q2	What is your gender?	Female	2	Attribute
Q7	Do you consider public transport to be affordable?	Not Expensive	3	Condition
Q7	Do you consider public transport to be affordable?	Expensive	4	Condition

D3. Rewriting to ADICO components (LLM API Example)

Table D.4: Example Data of Question and Response Options Rewritten as IG Components

Question_Response_ID	Question_Text	IG_Component
Q1_3	What is your age?	if they are 50 or younger
Q1_4	What is your age?	if they are over 50
Q2_1	What is your gender?	if they are male
Q2_2	What is your gender?	if they are female
Q2_3	What is your gender?	if they are other (example failed conversion)
Q4_1	Do you use public transportation frequently?	use public transportation frequently
Q4_2	Do you use public transportation frequently?	not use public transportation frequently
Q5_1	Did you vote in the last election?	achieve their right to vote
Q5_2	Did you vote in the last election?	decline their right to vote
Q6_1	Do you think recycling is easy?	if they think recycling is easy
Q6_2	Do you think recycling is easy?	if they think recycling is not easy
Q7_3	Do you consider public transport to be affordable?	if they consider public transport to be not expensive
Q7_4	Do you consider public transport to be affordable?	if they consider public transport to be expensive
Q8_1	Do you follow political news on Social media?	if they follow political news on Social media
Q8_2	Do you follow political news on Social media?	if they do not follow political news on Social media

Q10_1	Do your peers expect you to use public transportation?	if they feel their peers expect them to use public transportation
Q10_2	Do your peers expect you to use public transportation?	if they feel their peers do not expect them to use public transportation

D4. Create Parallel Set Diagrams

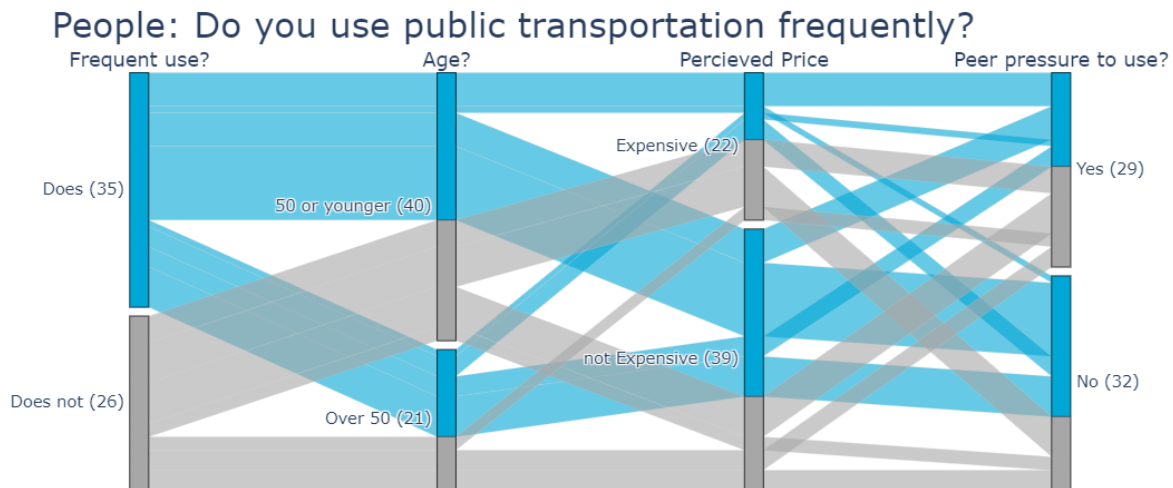


Figure D.1: Example Parallel Set Diagram from Survey Responses

D5. Extracting Statements Using Decision Trees

Q4: Do you use public transportation frequently?

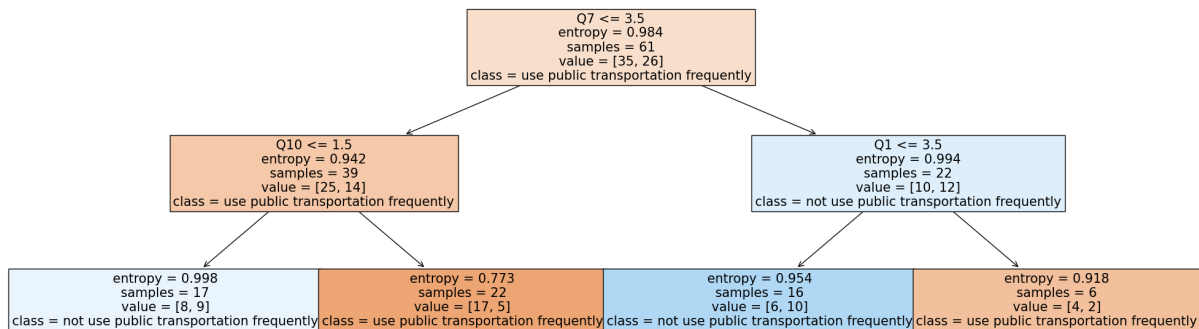


Figure D.2: Example Tree Diagram from Survey Responses (Iteration 1 of Using Public Transport)

Q4: Do you use public transportation frequently?

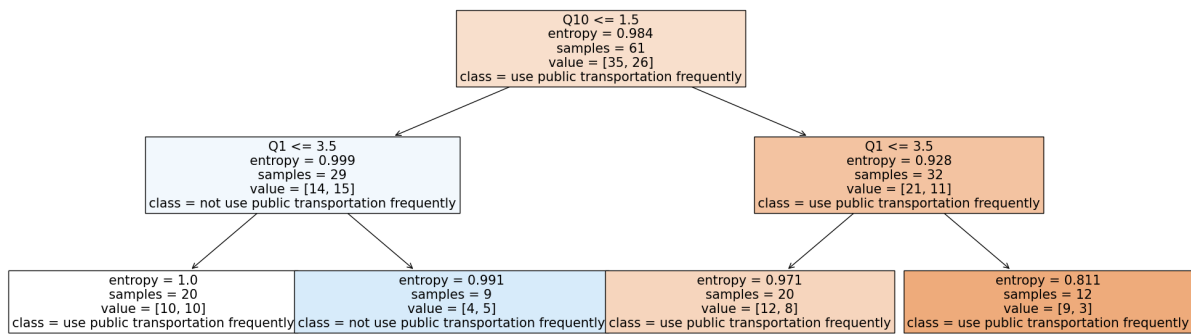


Figure D.3: Example Tree Diagram from Survey Responses (Iteration 2 of Using Public Transport)

Table D.5: Example Data of Extracted Informal Rule Statements and Data from Decision Trees

Attribute	Aim resp	Total	Aim %	C1_resp	C1 count	C1 Aim%	C2_resp	C2 count	C2 Aim%	Entropy
People	use public transportation frequently	61	0.57	if they consider public transport to be not expensive	39	0.64	if they feel their peers do not expect them to use public transportation	22	0.77	0.77
People	use public transportation frequently	61	0.57	if they feel their peers do not expect them to use public transportation	32	0.66	if they are over 50	12	0.75	0.81

D6. Visualising Extracted Statements

Sankey plot representation of an IG statement extracted by decision tree algorithm:
use public transportation frequently

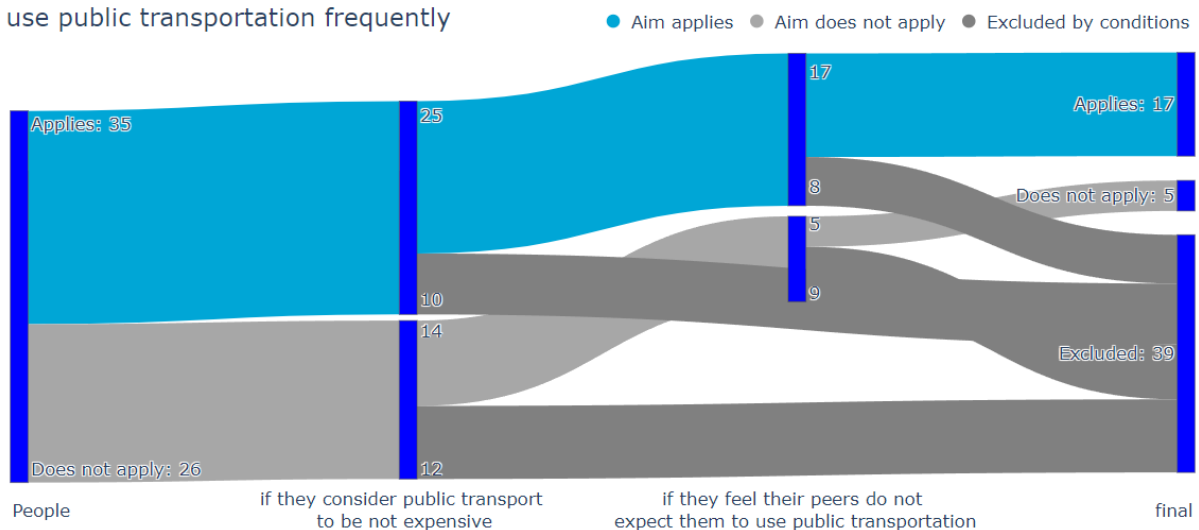


Figure D.4: Example Visualisation of Statement 1 extracted from Decision Tree Figure D.1

Sankey plot representation of an IG statement extracted by decision tree algorithm:
use public transportation frequently

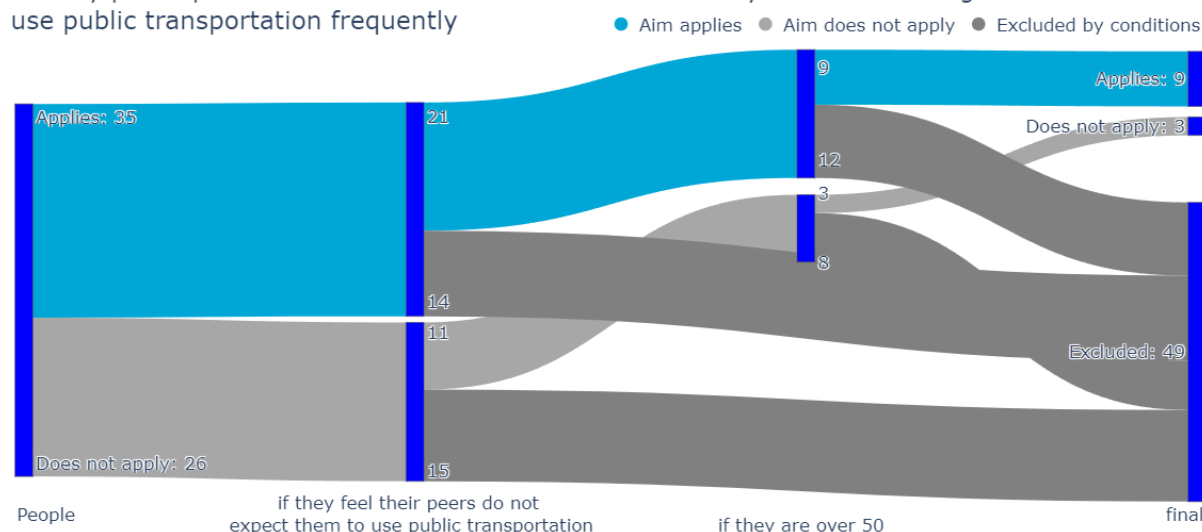


Figure D.5: Example Visualisation of Statement 1 extracted from Decision Tree Figure D.2

E. Limitations of Institutional Grammar and Mitigation Strategies

While Institutional Grammar (IG) provides a structured approach to analysing informal rules, it is important to acknowledge its limitations. This appendix outlines key challenges in applying IG to survey data and suggests mitigation strategies. These insights are crucial for researchers considering the use of the SIRE method or similar approaches in institutional analysis.

1. Rigidity in Structure:

Limitation: IG's structured syntax (ADICO) may be too rigid for capturing the nuances of informal rules, which are often fluid and context-dependent.

Impact: This could lead to oversimplification of complex social norms or behaviours.

Mitigation: Incorporate qualitative analysis alongside IG to capture nuances that don't fit neatly into the ADICO structure.

2. Difficulty with Implicit Rules:

Limitation: IG may struggle to capture implicit or unspoken rules that are not easily articulated in survey responses.

Impact: Important informal rules might be missed or misinterpreted.

Mitigation: Use complementary methods like ethnographic research or in-depth interviews to uncover implicit rules.

3. Over-reliance on Explicit Statements:

Limitation: IG typically relies on explicit statements, which might not always be present in survey data.

Impact: This could result in missing important informal rules that are implied rather than stated.

Mitigation: Develop methods to infer rules from patterns of responses, not just explicit statements.

4. **Potential for Oversimplification:**

Limitation: The structured nature of IG might lead to oversimplification of complex social phenomena.

Impact: This could result in a loss of nuance or misrepresentation of complex social dynamics.

Mitigation: Use IG in conjunction with other analytical methods to provide a more comprehensive understanding.

5. **Difficulty with Dynamic Rules:**

Limitation: IG might struggle to capture rules that change over time or in different contexts.

Impact: This could lead to a static representation of what might be dynamic informal rules.

Mitigation: Incorporate longitudinal analysis or scenario-based approaches to capture rule dynamics.

In conclusion, while IG offers a valuable framework for analysing informal rules, researchers must be aware of these limitations when applying it to survey data. By employing the suggested mitigation strategies and combining IG with complementary methods, researchers can enhance the robustness and comprehensiveness of their analyses. The SIRE method, developed in this thesis, represents an initial step in addressing some of these limitations, particularly in the context of quantitative survey data. Future research should continue to refine and adapt IG techniques to better capture the complexity of informal rules in diverse institutional settings.