



Attribute Non-Attendance in Discrete Choice Experiments:

Evidence from Swedish Investors and Farmers

Shahmeer Akhtar

Master thesis in Economics, A2E • 30 credits

Swedish University of Agricultural Sciences, SLU

Faculty of Natural Resources and Agricultural Sciences/Department of Economics

Agricultural, Food and Environmental Policy Analysis (AFEPA)

Degree project/SLU, Department of Economics 1756 • ISSN 1401-4084

Uppsala, Sweden 2026



Attribute Non-Attendance in Discrete Choice Experiments: Evidence from Swedish Investors and Farmers

Shahmeer Akhtar

Supervisor:	Jens Rommel, Swedish University of Agricultural Sciences, Department of Economics
Assistant supervisor:	Uliana Gottlieb, Swedish University of Agricultural Sciences, Department of Economics
Examiner:	Shon Ferguson, Swedish University of Agricultural Sciences, Department of Economics
Credits:	30 credits
Level:	Second cycle, A2E
Course title:	Master thesis in Economics, A2E
Course code:	EX0905
Programme/education:	Agricultural, Food and Environmental Policy Analysis (AFEPA)
Course coordinating dept:	Department of Economics
Place of publication:	Uppsala, Sweden
Year of publication:	2026
Title of series:	Degree project/SLU, Department of Economics
Part number:	1756
ISSN:	1401-4084
Keywords:	Attribute Non-Attendance (ANA), Discrete Choice Experiments (DCE), Latent Class Models, Preference Heterogeneity, Stated vs Inferred ANA, Investor Behaviour, Farmer Decision-Making

Swedish University of Agricultural Sciences

Faculty of Natural Resources and Agricultural Sciences
Department of Economics

Abstract

Attribute non-attendance (ANA) is a common behavioural phenomenon in discrete choice experiments (DCEs) and can lead to biased estimates if not properly accounted for. This study investigates ANA in two applied contexts using data from farmer and investor choice experiments. The analysis combines descriptive indicators of stated ANA with latent class models to infer non-attendance behaviour from observed choices. In addition, the relationship between stated and inferred ANA is examined to assess the consistency between self-reported and behaviourally derived measures. The results provide consistent evidence of attribute non-attendance in both datasets, with variation across attributes and contexts. For farmers, financial incentives play a central role in decision-making, while for investors, risk is a key determinant of choices. The comparison between stated and inferred ANA reveals limited alignment, suggesting that self-reported behaviour may not fully capture underlying decision processes. Overall, the findings highlight the importance of accounting for ANA in DCE applications and provide insights into how attribute processing differs across contexts.

Keywords: Attribute Non-Attendance (ANA), Discrete Choice Experiments (DCE), Latent Class Models, Preference Heterogeneity, Stated vs Inferred ANA, Investor Behaviour, Farmer Decision-Making

Table of contents

List of tables	6
List of figures	7
Abbreviations	8
1. Introduction	9
2. Literature review	11
2.1 Discrete choice experiments in environmental and agricultural applications	11
2.2 Attribute non-attendance in discrete choice experiments	11
2.3 Stated and inferred ANA	12
2.4 Latent class models and inferred ANA	12
3. Data	15
3.1 Data sources	15
3.2 Discrete choice experiment design	15
3.2.1 Farmer dataset	16
3.2.2 Investor dataset	17
3.3 Measurement of attribute non-attendance (ANA)	18
4. Methodology	19
4.1 Random utility framework	19
4.2 Latent class model	19
4.3 Modelling attribute non-attendance	20
4.4 Class membership specification	21
4.5 Stated and inferred ANA	22
4.6 Empirical strategy and estimation	22
4.7 Extension: multi-class specification	23
5. Results	24
5.1 Investor data	24
5.1.1 Stated attribute non-attendance	24
5.1.2 Baseline latent class results	25
5.1.3 Baseline versus zero-constrained models	25
5.1.4 Class membership and covariates	26
5.1.5 Stated versus inferred ANA	27
5.1.6 Summary of main findings	28
5.1.7 Robustness check: restricted stated ANA sample	28
5.1.8 Extension: Multi-class (n+1) specification	29
5.2 Farmer data	30
5.2.1 Stated attribute non-attendance	30
5.2.2 Baseline latent class results	31
5.2.3 Baseline versus zero-constrained models	31
5.2.4 Class membership and covariates	32
5.2.5 Stated versus inferred ANA	32
5.2.6 Summary of main findings	33
5.2.7 Robustness check: restricted stated ANA sample	34
6. Discussion and Conclusion	35
References	39
Popular science summary	42
Appendix	43

List of tables

Table 1. Selected attribute non-attendance studies relevant to this thesis.....	13
Table 2. Attributes and levels in the farmer experiment.....	16
Table 3. Attributes and levels in the investor experiment	17
Table 4. Estimated class shares (baseline latent class models, investor dataset)	25
Table 5. Likelihood ratio test results (investor dataset).....	26
Table 6. Average probability of Class B by stated ANA (investor dataset).....	27
Table 7. Correlation between stated ANA and Class B probability (investor dataset).....	28
Table 8. Estimated class shares (baseline latent class models, farmer dataset)	31
Table 9. Likelihood ratio test results (farmer dataset).....	31
Table 10. Average probability of Class B by stated ANA (farmer dataset)	32
Table 11. Correlation between stated ANA and Class B probability (farmer dataset)	33
Table 12. Model fit statistics (two-class latent class models, investor dataset).....	45
Table 13. Model fit statistics with restricted stated ANA (investor dataset, N = 730)	45
Table 14. Estimated class probabilities (n+1 latent class model, investor dataset).....	45
Table 15. Effect of T60 treatment on stated ANA (investor dataset)	46
Table 16. Class membership estimates (investor dataset)	46
Table 17. Class membership estimates (farmer dataset)	46
Table 18. Model fit statistics (two-class latent class models, farmer dataset)	47
Table 19. Model fit statistics with restricted stated ANA (farmer dataset, N = 279).....	47

List of figures

Figure 1. Stated ANA shares across attributes in the investor dataset.	24
Figure 2. Stated ANA shares across attributes in the farmer dataset.	30
Figure 3. Example of a choice task from the farmer dataset	43
Figure 4. Example of a choice task from the investor dataset	43
Figure 5. Distribution of ANA Class Probabilities by Attribute (Investor data)	44
Figure 6. Distribution of ANA Class Probabilities by Attribute (Farmer data)	44

Abbreviations

ANA	Attribute Non-Attendance
DCE	Discrete Choice Experiment
EU	European Union
LC	Latent Class
MLE	Maximum Likelihood Estimation
MNL	Multinomial Logit
RUM	Random Utility Model
WTP	Willingness to Pay
AIC	Akaike Information Criterion
BIC	Bayesian Information Criterion
LR	Likelihood Ratio

1. Introduction

Discrete choice experiments (DCEs) are widely used in environmental and agricultural policy analysis to value non-market goods for which no observable market prices exist (Alpizar et al., 2001; Hoyos, 2010). Many environmental outcomes, such as biodiversity conservation, reduced pollution, and ecosystem services, are central to policy design but cannot be directly traded in markets. DCEs provide a structured framework to elicit preferences by asking respondents to choose between hypothetical alternatives described by attributes and associated costs. Previous studies have highlighted the suitability of DCEs for analysing agri-environmental measures and climate-related policies, particularly in situations where revealed preference data are limited or unavailable (Hoyos, 2010; Rakotonarivo et al., 2016). Recent review studies also show that DCEs are widely applied in environmental valuation and in the design of agri-environmental contracts (Schulze et al., 2024).

A key assumption underlying standard DCE models is that respondents consider all attributes presented in the choice tasks. However, this assumption is often violated. Previous research suggests that respondents may simplify complex decisions by ignoring one or more attributes when evaluating alternatives, a behaviour commonly referred to as attribute non-attendance (ANA) (Hensher et al., 2005; Scarpa et al., 2009). If ANA is not accounted for, estimated preferences and welfare measures may be biased, which can lead to misleading policy conclusions. Gonçalves et al. (2022) provide theoretical and empirical evidence that attribute non-attendance can influence decision-making and results in choice experiments. Earlier studies also show that respondents may ignore specific attributes, such as cost, and that such behaviour can be related to information load and task complexity.

An important distinction in the ANA literature is between stated and inferred ANA. Stated ANA refers to self-reported information on which attributes respondents claim to have ignored, whereas inferred ANA is derived from observed choices using econometric models. These two approaches do not necessarily coincide, which raises questions about how well self-reports reflect actual decision processes. Given these inconsistencies, several studies have emphasised the importance of examining the relationship between stated and inferred ANA when interpreting DCE results (Campbell et al., 2011; Hensher & Greene, 2010).

Despite this growing literature, relatively little evidence exists on how the relationship between stated and inferred ANA varies across fundamentally different decision contexts. This thesis aims to analyse attribute non-attendance in discrete choice experiments using latent class models, with a focus on comparing stated and inferred ANA across two applied contexts. The analysis is based on farmer and investor choice experiment data, which differ in both their decision environment and attribute structure. Specifically, the thesis examines whether ANA is present across attributes and contexts, and whether stated and inferred

measures provide consistent insights into respondents' behaviour. While previous studies have documented the presence of attribute non-attendance across various contexts, less is known about how ANA compares across fundamentally different decision environments and to what extent stated measures align with behaviour inferred from observed choices. This gap motivates the empirical analysis conducted in this thesis. The study contributes to the literature in two ways. First, it provides a comparative analysis of ANA across two distinct application domains, allowing differences in attribute processing to be studied across contexts. Second, it combines descriptive evidence on stated ANA with latent class models that infer ANA from observed choices. In doing so, the thesis contributes to the broader literature on how respondents process information in DCEs and how ANA should be incorporated into empirical analysis.

To guide the analysis, this study addresses the following research questions. First, to what extent is attribute non-attendance present across attributes and contexts? Second, is attribute non-attendance uniform, or does it vary across attributes? Third, does accounting for attribute non-attendance improve model fit in latent class models? Fourth, to what extent do stated and inferred measures of attribute non-attendance align? Finally, how does attribute processing differ between investor and farmer decision contexts?

2. Literature review

2.1 Discrete choice experiments in environmental and agricultural applications

DCEs have become an important method in environmental valuation because they allow researchers to analyse preferences for non-market goods and policy trade-offs that cannot be observed in actual markets. Hoyos (2010) reviews the state of the art in environmental valuation with DCEs and shows their increasing relevance in environmental decision-making. Similarly, Rakotonarivo et al. (2016) systematically review environmental DCE applications and show that DCEs are widely used, while also noting that evidence on their reliability and validity is mixed.

In agricultural policy research, DCEs are also increasingly used to analyse farmers' preferences for contract design, incentives, and environmental measures. A recent systematic review by Schulze et al. (2024) examines 127 DCE studies on farmers' ex ante preferences for agri-environmental contracts, showing how widely the method is used in this area and the importance of contract attributes in shaping participation decisions. Recent work by Villanueva et al. (2026) also demonstrates the increasing use of DCEs in agri-environmental policy research, especially in the context of land managers' participation in environmental and sustainability programs. This literature provides strong support for using DCEs in both the farmer and investor contexts analysed in this thesis, where preferences over policy-relevant attributes are central.

2.2 Attribute non-attendance in discrete choice experiments

Although standard DCE models assume that respondents evaluate all attributes, this assumption is often unrealistic. A large literature shows that respondents may ignore some attributes in order to simplify the choice task. Gonçalves et al. (2022) discuss the role of attribute non-attendance in consumer decision-making and show that ignoring attributes can affect empirical results. The review also highlights that ignoring attributes can substantially affect estimated preferences and welfare measures, particularly when monetary attributes are included in the choice tasks. Early contributions by Hensher et al. (2005) show that attribute non-attendance can have important implications for willingness-to-pay estimates. Hensher (2006) further demonstrates that the extent to which respondents ignore attributes is related to the complexity of the choice task and the information load they face. These studies suggest that ANA may reflect a deliberate simplification strategy rather than random error. In environmental applications, Scarpa et al. (2009) and Carlsson et al. (2010) also show that ignoring attributes can affect valuation results and policy interpretation. Decision-making in complex

environments is often constrained by cognitive limitations. As noted by Simon (1955):

“the capacity of the human mind for formulating and solving complex problems is very small compared with the size of the problems whose solution is required...”

This provides a theoretical foundation for interpreting attribute non-attendance as a simplifying strategy in discrete choice experiments.

2.3 Stated and inferred ANA

The literature commonly distinguishes between stated and inferred attribute non-attendance (ANA). Stated ANA is based on respondents’ own reports about which attributes they ignored after completing the choice tasks, whereas inferred ANA is obtained from observed choices using econometric models that identify low or zero sensitivity to particular attributes (Alemu et al., 2013; Hensher et al., 2005; Scarpa et al., 2009). This distinction is important because self-reported non-attendance may not fully reflect the actual behavioural rules used during decision-making. Comparing stated and inferred ANA has therefore become an important part of empirical research. Several studies find limited correspondence between the two measures, suggesting that respondents are only partially aware of their own simplification strategies or that self-reported measures capture these imperfectly (Campbell et al., 2011; Hensher & Greene, 2010). As a result, relying solely on stated ANA may lead to misleading conclusions about attribute processing. A joint consideration of both stated and inferred measures is therefore useful for assessing the extent to which reported behaviour aligns with observed choice patterns. Furthermore, awareness or visual attention to an attribute does not necessarily imply that it is considered in the decision process, as individuals may process information without incorporating it into their final choice. This suggests that different approaches to measuring ANA capture distinct aspects of the decision process. Consequently, combining stated and inferred ANA provides a more comprehensive understanding of how individuals process information in discrete choice experiments. Recent research also shows that self-reported consideration may partly reflect underlying preferences rather than actual decision rules (Edenbrandt et al., 2022).

2.4 Latent class models and inferred ANA

A prominent approach for modelling inferred ANA is the latent class model. In this framework, respondents are allocated probabilistically to different classes, where each class represents a distinct pattern of preferences or attribute-processing behaviour. Hensher & Greene (2010) and Campbell et al. (2011) show that latent class models are particularly suitable for ANA analysis because they can represent classes in which one or more attributes are not attended to. In many ANA applications, non-attendance is represented by constraining the coefficient of an attribute to zero within a given class. In this case, the attribute does not contribute to utility for respondents in that class. This zero-constrained latent class logic is central to the empirical strategy used in this thesis, as it provides a direct

way to compare standard latent class specifications with models that explicitly impose non-attendance. Latent class models can also be extended by allowing class membership probabilities to depend on respondent characteristics. Information-processing behaviour is unlikely to be homogeneous across individuals and may vary with factors such as experience, cognitive burden, and other characteristics (Campbell et al., 2011; Hensher & Greene, 2010). Including covariates in the class membership function therefore helps explain heterogeneity in the probability of belonging to different latent classes, including those associated with attribute non-attendance. In this thesis, demographic and behavioural covariates are included in the class allocation stage to explain variation in class membership probabilities. These covariates do not directly affect utility, but rather capture systematic differences in the likelihood that individuals follow different attribute-processing strategies.

Gonçalves et al. (2022) provide a systematic review of attribute non-attendance in discrete choice experiments and distinguish between stated and inferred approaches, as well as between different econometric strategies used to account for non-attendance. Their review shows that ANA has been studied across several fields, including transport, environmental economics, health, and food/agricultural applications. It also highlights that inferred ANA is often modelled using latent class specifications in which ignored attributes are assigned zero utility weights. Accounting for ANA can affect model fit and willingness-to-pay estimates, although the direction and size of these effects vary across studies and contexts. A further issue in this literature is that inferred non-attendance may be difficult to separate from weak preferences or broader preference heterogeneity. Respondents may appear to ignore an attribute because they place little weight on it, rather than because they follow a strict non-attendance rule. This distinction is relevant for the present thesis, where stated and inferred ANA are compared across two different decision contexts.

Table 1 summarises selected studies from the ANA literature that are most relevant to the empirical approach used in this thesis. The table is not intended as a full systematic review, since such reviews already exist, but rather as a compact overview of studies that inform the use of stated ANA, inferred ANA, latent class models, and zero-constrained specifications.

Table 1. Selected attribute non-attendance studies relevant to this thesis

Study	Context	ANA type	Model / approach	Relevance for this thesis
(Hensher et al., 2005)	Transport	Stated ANA	Mixed logit with ignored attributes constrained to zero	Early example showing that ignoring attributes can affect willingness-to-pay estimates
(Scarpa et al., 2009)	Environmental valuation	Inferred ANA	Equality-constrained latent class model	Key reference for zero-constrained class-based ANA
(Hensher & Greene, 2010)	Transport	Inferred ANA	Latent class model	Shows how latent classes can represent different attribute-processing strategies

(Campbell et al., 2011)	Environmental valuation	Inferred ANA	Latent class specification	Shows that ANA can be modelled through classes with different attendance patterns
(Alemu et al., 2013)	Environmental valuation	Stated ANA	Follow-up ANA questions including reasons for non-attendance	Important because stated reasons may reflect preferences, not only simplification
(Caputo et al., 2018)	Food choice	Stated and inferred ANA	Comparison of serial, choice-task, and inferred ANA methods	Relevant for comparing different ways of measuring ANA
(Alho, 2017)	Investment choices	Stated and inferred ANA	Equality-constrained latent class model	Directly relevant because it compares stated and inferred use of investment attributes
(Yangui et al., 2025)	Agricultural economics	Stated ANA	Comparison of ANA across different elicitation formats	Recent evidence on how ANA measurement varies across elicitation formats

Notes: The table presents selected studies that inform the empirical approach used in this thesis. The focus is on stated and inferred ANA, latent class specifications, and zero-constrained treatments of non-attendance. It is not intended as a full systematic review. Studies such as Alho (2017), Caputo et al. (2018), and Yangui et al. (2025) further show that stated and inferred ANA can differ depending on context and elicitation format.

Overall, the selected literature shows that ANA can be measured and modelled in different ways, and that accounting for non-attendance can affect model fit, willingness-to-pay estimates, and the interpretation of preferences. While ANA has been studied in several applied contexts, there is limited evidence comparing stated and inferred ANA across decision contexts as different as agricultural production decisions and sustainable investment choices. This comparative setting therefore represents an important empirical contribution of the study. The latent class modelling approach outlined above provides the foundation for the empirical strategy presented in the following sections.

3. Data

3.1 Data sources

This study uses two datasets based on discrete choice experiments (DCEs), each designed to analyse decision-making behaviour in different applied contexts.

The first dataset originates from a discrete choice experiment conducted among Swedish dairy farmers to investigate preferences for methane-reducing feed additives (Gottlieb & Rommel, 2026). The survey was implemented by the Swedish University of Agricultural Sciences (SLU) in collaboration with Next Research & Consulting using contact information obtained through Statistics Sweden (SCB). The final sample consists of 320 respondents and was collected to examine how advisory provision, methane reduction outcomes, financial incentives, and climate-related information influence farmers' adoption decisions. Respondents completed six discrete choice tasks involving alternative feed additive adoption programmes and an opt-out option. Further details regarding the survey design and implementation are provided by Gottlieb and Rommel (2026).

The second dataset originates from an online discrete choice experiment conducted among Swedish retail investors in June 2024 to investigate preferences for environmentally oriented equity funds and the role of investment stake size. The survey was administered through a professional market research company and the final sample consists of 804 respondents. The study was designed to examine preferences for sustainable investment funds characterised by environmental objectives, management strategies, environmental performance information, geographical scope of impact, risk levels, and expected returns. To ensure relevance and data quality, respondents working professionally with investments, lacking experience with investment funds, not actively investing in funds, or failing attention checks were excluded from the final sample.

Both datasets were collected in Sweden and represent two distinct decision-making contexts: agricultural production decisions and financial investment choices. Both datasets are structured in panel (long) format, where each respondent completes multiple choice tasks. As a result, the number of observations exceeds the number of individuals. For instance, the investor dataset consists of 804 respondents, each completing multiple choice tasks, resulting in approximately 4,800 choice observations.

3.2 Discrete choice experiment design

Both datasets are based on discrete choice experiments grounded in random utility theory. In each experiment, respondents are asked to choose between alternatives described by multiple attributes, allowing the estimation of preferences and trade-offs.

3.2.1 Farmer dataset

In the farmer dataset, respondents completed six repeated choice tasks involving two hypothetical feed additive programmes and an opt-out alternative representing continuation of current farming practices. The choice tasks are designed to reflect realistic decisions farmers face when considering the adoption of new practices or technologies, including trade-offs between economic incentives and implementation conditions.

The attributes and levels used in the experiment are presented in Table 2.

Table 2. Attributes and levels in the farmer experiment

Attribute	Levels
Who gives advice	No advice offered (reference) Farm advisor Veterinarian Experienced farmer
How often advice is given	No advice offered (reference) Once per year Twice per year Three times per year
Reduction of methane emissions	10% reduction 20% reduction 30% reduction 40% reduction
Net bonus per kg of milk	0 öre (reference) 5 öre 10 öre 15 öre 20 öre 25 öre 30 öre 35 öre

A sample choice card is shown in Appendix (Figure 3). The example illustrates how respondents choose between two alternatives (Feed additive A and B) and an opt-out option (no adoption). The opt-out option represents maintaining current practices. Each alternative is described by attributes such as advisory source, frequency of advice, environmental impact (methane reduction), and economic incentives.

In addition, the farmer experiment includes a between-group information treatment. Respondents were randomly assigned to either a control group or a treatment group that received additional information on climate policy goals related to methane reduction. This allows the analysis of how information provision influences preferences and decision-making. Prior to the choice tasks, respondents were provided with clear explanations of all attributes and instructions to ensure understanding of the decision context. The survey included background and screening questions, followed by the choice experiment and

subsequent follow-up questions, including those used to construct attribute non-attendance indicators.

3.2.2 Investor dataset

In the investor dataset, respondents are presented with repeated choice situations in which they choose between two investment alternatives and an opt-out option representing no investment. This corresponds to maintaining their current investment allocation. Each respondent completed six choice tasks, consistent with standard practices in discrete choice experiments to balance statistical efficiency and respondent fatigue. The experimental design was constructed using a Bayesian D-efficient design approach, which optimises the statistical efficiency of parameter estimation based on prior information about expected preferences (Hensher et al., 2015). The full design consisted of 24 choice situations divided into four blocks, with respondents randomly assigned to one block and completing six choice tasks. This approach reduced respondent burden while maintaining statistical efficiency. A sample choice card is shown in Appendix (Figure 4). The attributes and levels used in the experiment are presented in Table 3.

Table 3. Attributes and levels in the investor experiment

Attribute	Levels
Environmental objective	Transition to circular economy Pollution prevention and control Protection and restoration of biodiversity and ecosystems Climate change mitigation and adaptation
Management strategy	Exclude Select Influence
Region of operations	Global Sweden Europe
Environmental performance information	Assured by an auditor Environmental score issued by a third party
Expected annual return	6% 9% 12% 15%
Risk classification	3 (lower medium) 4 (medium) 5 (upper medium) 6 (second highest)

A key feature of the design is a stake treatment, where respondents were randomly assigned to invest either a lower or a higher share of their monthly investment savings (20% or 60%). This enables the analysis of behaviour under different levels of financial commitment and allows examination of whether preferences for sustainable investment attributes vary with the magnitude of the investment decision.

Respondents were provided with detailed explanations of all attributes prior to the choice tasks. To ensure understanding, comprehension checks in the form of true/false questions were included. Respondents who answered incorrectly were required to review the attribute descriptions before repeating the comprehension test, ensuring adequate understanding of the decision context before proceeding to the choice tasks.

Across both experiments, the attributes are designed to reflect realistic and policy-relevant decision environments. In the farmer dataset, the attributes focus on advisory provision, environmental outcomes in terms of methane emission reduction, and economic incentives linked to adoption decisions. In contrast, the investor dataset captures financial and sustainability-related dimensions of investment choices, including expected returns, risk levels, and environmental characteristics of funds. Taken together, the two experimental designs allow for the analysis of trade-offs between economic, environmental, and informational factors across two distinct but comparable decision-making contexts.

3.3 Measurement of attribute non-attendance (ANA)

Both datasets include information that allows the analysis of attribute non-attendance (ANA). In addition to the choice tasks, respondents were asked follow-up questions indicating whether they ignored specific attributes when making their decisions. These responses are used to construct stated ANA indicators for each attribute. Specifically, respondents were asked to report whether they paid attention to each attribute when completing the choice tasks. Binary indicators are constructed, taking the value of one if an attribute is reported as ignored and zero otherwise. Respondents could also indicate that no attributes were ignored. This approach follows the stated ANA literature, which uses self-reported measures to capture attribute-processing behaviour and decision heuristics (Alemu et al., 2013). Importantly, the reasons for stated non-attendance may differ across respondents. Alemu et al. (2013) highlight that respondents may report non-attendance for different reasons, and that reported non-attendance does not necessarily reflect simplification behaviour. For example, respondents may ignore an attribute because it is not relevant to their decision-making, rather than due to cognitive constraints. Therefore, stated ANA is interpreted cautiously, as it may capture both true preference heterogeneity and decision heuristics. In addition to stated ANA, the panel structure of the data allows for the estimation of inferred ANA using econometric models. In particular, latent class models are employed, where non-attendance is represented by zero-constrained parameters for specific attributes within a given class. This enables the identification of behavioural patterns consistent with attribute non-attendance based on observed choices.

4. Methodology

The modelling strategy closely follows established approaches in the ANA literature, while extending them to a comparative multi-context setting.

4.1 Random utility framework

The empirical analysis is based on the random utility framework commonly used in discrete choice experiments (McFadden, 1974; Train, 2009). For individual i , alternative j , and choice situation t , utility is defined as:

$$U_{ijt} = V_{ijt} + \varepsilon_{ijt} \quad (1)$$

where V_{ijt} represents the observable component of utility and ε_{ijt} is an unobserved stochastic term. Individuals are assumed to choose the alternative that yields the highest utility. The deterministic component of utility is specified as a linear function of observed attributes:

$$V_{ijt} = \beta' X_{ijt} \quad (2)$$

where X_{ijt} is a vector of attribute levels and β' is a vector of coefficients to be estimated. Both datasets are treated as panel data, as each respondent completes multiple choice tasks.

4.2 Latent class model

To account for preference heterogeneity and differences in decision-making behaviour, a latent class framework is employed (Greene & Hensher, 2003; Train, 2009). In this approach, individuals are probabilistically assigned to a finite number of latent classes, each characterised by a distinct set of preference parameters.

Conditional on class membership, choices follow a multinomial logit structure (McFadden, 1974). For class s , the probability that individual i chooses alternative j in choice situation t is:

$$P_{ijt|s} = \frac{\exp(V_{ijt|s})}{\sum_{k=1}^J \exp(V_{ikt|s})} \quad (3)$$

where J denotes the number of alternatives in the choice set, and k indexes the alternatives included in the choice set. The likelihood for each individual is then constructed by multiplying probabilities across all observed choice tasks, reflecting the panel structure of the data. The unconditional likelihood for individual i is given by:

$$P_i = \sum_{s=1}^S \pi_{is} \prod_{t=1}^{T_i} P_{it|s} \quad (4)$$

Where s denotes the latent class, π_{is} is the probability that individual i belongs to class s , $P_{it|s}$ is the class-specific probability of the observed choice and T_i is the number of choice situations faced by individual i . The baseline specification considers two latent classes: one class with full attribute attendance (Class A) and another class characterised by attribute non-attendance (Class B).

4.3 Modelling attribute non-attendance

Attribute non-attendance (ANA) is incorporated within the latent class framework by imposing zero constraints on selected coefficients. Following the latent class approach discussed in Section 2.4, non-attendance is operationalised by constraining attribute coefficients to zero within specific latent classes (Campbell et al., 2011; Scarpa et al., 2009).

Specifically, for attribute k in the non-attendance class:

$$\beta_{k,B} = 0 \quad (5)$$

This restriction implies that the corresponding attribute does not contribute to utility for individuals belonging to that class.

For each attribute group, two model specifications are estimated:

- a baseline model, where all coefficients are freely estimated in both classes,
- a zero-constrained model, where coefficients corresponding to a given attribute are fixed to zero in the second class.

In the zero-constrained specifications, Class A represents the full-attendance class, while Class B represents the potential non-attendance class for the respective attribute. The comparison between these specifications allows the identification of behavioural patterns consistent with attribute non-attendance.

The attribute groups considered in the analysis correspond to those described in Section 3.2. While the specific attributes differ across datasets due to differences in experimental design, the modelling approach remains consistent, with zero constraints applied to the relevant coefficients depending on the attribute under consideration.

4.4 Class membership specification

Class membership probabilities are modelled as a function of individual-specific characteristics (Hensher & Greene, 2010). Let Class A be the reference class. The utility of belonging to Class B is specified as:

$$M_{iB} = \delta_B + \gamma_{ANA} ANA_i + \gamma' Z_i \quad (6)$$

where δ_B is a class-specific constant, ANA_i is the stated ANA indicator and Z_i is a vector of covariates.

The probability of belonging to Class B is given by the logistic form:

$$\pi_{iB} = \frac{\exp(M_{iB})}{1 + \exp(M_{iB})}, \pi_{iA} = 1 - \pi_{iB} \quad (7)$$

where π_{iB} denotes the probability that individual i belongs to Class B, and π_{iA} is the corresponding probability of belonging to Class A.

Covariates are included only in the class membership model and not in the utility specification. This allows observed characteristics to explain heterogeneity in attribute-processing behaviour rather than directly entering the utility function.

For the farmer dataset, covariates include demographic and behavioural variables such as gender, age, risk-taking, and trust. For the investor dataset, class membership probabilities are modelled as a function of stated ANA and a set of behavioural and demographic covariates, including anticipated regret, perceived attractiveness of the investment, gender, deliberative decision-making style, and trust in information sources. Covariates are standardised prior to estimation where used in the class membership model, while ANA variables are coded as binary indicators. This transformation improves numerical stability and facilitates comparison across variables measured on different scales.

The inclusion of covariates in the class membership function reflects the idea that attribute-processing behaviour is not homogeneous across individuals, but may vary with cognitive effort, experience, and engagement with the choice task (Hensher, 2006). However, the role of these variables is interpretative rather than causal, as the model identifies associations between observed characteristics and the probability of belonging to different latent classes.

To evaluate the role of attribute non-attendance, the baseline and zero-constrained models are compared using likelihood-based criteria (Train, 2009). The comparison assesses whether restricting selected coefficients to zero significantly affects model fit. The likelihood ratio statistic is defined as:

$$LR = -2(LL_R - LL_U) \quad (8)$$

where LL_R denotes the log-likelihood of the restricted (zero-constrained) model and LL_U denotes the log-likelihood of the unrestricted (baseline) model. Under

standard regularity conditions, the statistic follows a chi-squared distribution with degrees of freedom equal to the number of imposed restrictions.

This framework provides a consistent basis for evaluating whether zero-constrained specifications offer an improved representation of attribute-processing behaviour within the latent class setting.

4.5 Stated and inferred ANA

Consistent with the distinction outlined in section 2.3, the analysis distinguishes between stated and inferred ANA. Stated ANA is derived from survey responses indicating which attributes respondents report ignoring. Inferred ANA is captured through the latent class structure, where non-attendance is represented by zero-constrained coefficients.

To examine the relationship between these measures, class membership probabilities are computed for each individual (Hensher & Greene, 2010). These probabilities are compared with stated ANA indicators using correlation measures and summary comparisons to assess consistency between reported and inferred behaviour.

$$\rho = \text{corr}(\widehat{\pi}_{iB}, ANA_i) \quad (9)$$

This approach provides an indication of the extent to which self-reported attribute processing aligns with behaviour inferred from observed choice behaviour.

4.6 Empirical strategy and estimation

In addition to the latent class models, descriptive analyses are conducted to characterise patterns of stated attribute non-attendance. For each individual, ANA indicators are aggregated across choice tasks to construct respondent-level measures. These measures are used to examine the proportion of respondents ignoring each attribute, the number of attributes ignored, and summary statistics of relevant covariates.

In the investor dataset, additional regression-based analyses are conducted to explore associations between experimental treatments and ANA behaviour at both the aggregate and attribute-specific levels. These supplementary analyses provide additional context for interpreting the latent class results. All models are estimated by maximum likelihood using the Apollo package in R. The panel structure of the data is explicitly accounted for in the likelihood function.

The empirical strategy consists of estimating, for each attribute group, a baseline two-class latent class model and a corresponding zero-constrained specification, followed by likelihood-based comparisons of model fit. This estimation framework is applied consistently across both datasets.

4.7 Extension: multi-class specification

As an extension to the baseline framework, a more flexible latent class specification is explored for the investor dataset. In this case, an $n + 1$ class model is considered, where one class represents full attribute attendance and the remaining classes capture different patterns of attribute non-attendance across attributes. This extension is motivated by the possibility that attribute non-attendance may not be adequately captured by a single ANA class, particularly in datasets with larger sample sizes and more complex attribute structures. It therefore allows for a richer representation of heterogeneity in attribute processing. Due to the increased number of parameters and model complexity, this specification is treated as an exploratory extension to the main analysis.

5. Results

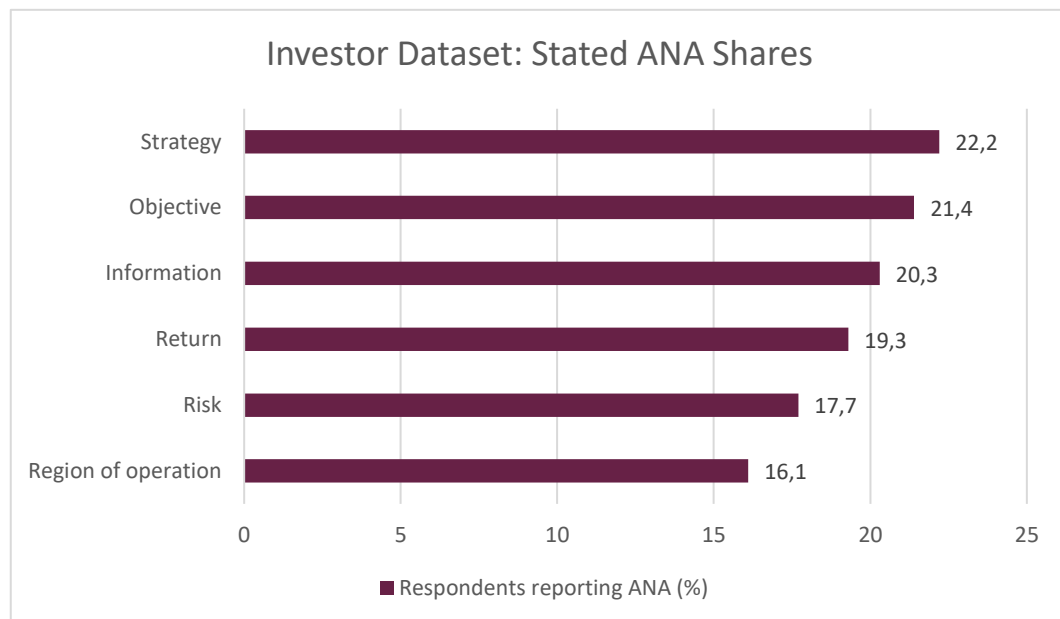
5.1 Investor data

5.1.1 Stated attribute non-attendance

We begin by examining stated attribute non-attendance (ANA) based on respondents' self-reported behaviour.

Overall, a meaningful share of respondents reported ignoring at least one attribute during the choice tasks. The proportion of respondents indicating non-attendance ranges from 16.1% to 22.2% across attributes. The highest shares are observed for management strategy (22.2%) and environmental objectives (21.4%), while region of operation (16.1%) and risk (17.7%) show comparatively lower levels. These values suggest that respondents do not ignore the choice task entirely, but instead simplify it by focusing on a subset of the available information. Figure 1 shows the share of respondents reporting attribute non-attendance for each attribute in the investor dataset.

Figure 1. Stated ANA shares across attributes in the investor dataset.



Note: Shares are reported separately for each attribute. Respondents could report non-attendance for multiple attributes; therefore, percentages do not sum to 100%.

At the respondent level, the average number of ignored attributes is approximately **1.17**. This indicates that simplification is common but not extreme. Most respondents ignore either no attributes or only one attribute, while ignoring multiple attributes is less frequent. Taken together, these descriptive results

provide initial evidence that ANA is present in the investor dataset and motivate the structural analysis that follows.

5.1.2 Baseline latent class results

To account for preference heterogeneity and potential attribute non-attendance (ANA), two-class latent class models were estimated separately for each attribute group. Across most baseline specifications, the class allocation results are highly stable. In the models for strategy, objective, return, information, and region of operation, the Class A (the dominant class) accounts for approximately 87–88% of the sample in each model, while the secondary class accounts for around 12–13%. This suggests that, for most attributes, the data support a relatively small segment of respondents who follow a different decision rule from the majority.

A clear exception emerges for the risk specification. In the baseline model for risk, the secondary class is considerably larger, representing 29.96% of respondents, compared with 70.04% in the dominant class. This indicates that heterogeneity related to risk is stronger than for the other attributes and may reflect differences in sensitivity to risk rather than straightforward attribute non-attendance.

Table 4. Estimated class shares (baseline latent class models, investor dataset)

Attribute	Class A	Class B
Strategy	87.21%	12.79%
Objective	87.60%	12.40%
Risk	70.04%	29.96%
Return	87.74%	12.26%
Information	87.66%	12.34%
Region of operation	87.58%	12.42%

The reported values represent the estimated average class membership probabilities (class shares) from each attribute-specific baseline model.

These baseline results indicate that respondents are not homogeneous in how they process the choice tasks. However, differences between classes in the baseline model do not by themselves imply attribute non-attendance. To evaluate whether the secondary class can be interpreted as a non-attendance class, the baseline models are compared with zero-constrained specifications. This distinction is important because the secondary class identified in the baseline model can only be interpreted as an ANA class if the zero restrictions are supported by the data.

5.1.3 Baseline versus zero-constrained models

The main test for inferred ANA compares each baseline model with a corresponding zero-constrained model, in which the coefficient of the relevant attribute is fixed to zero in Class B. If imposing the zero restriction does not significantly worsen model fit, the data are considered consistent with ANA for that attribute. If model fit deteriorates significantly, the restriction is not supported.

The likelihood ratio results show a consistent pattern. For strategy, objective, return, information, and region of operation, the zero-constrained specifications do not significantly worsen model fit. In contrast, for risk, the zero restriction leads to a statistically significant deterioration in fit ($LR = 23.45$, $p < 0.001$). This suggests that risk is not plausibly ignored in the same way as the other attributes.

Table 5. Likelihood ratio test results (investor dataset)

Attribute	LR statistic	df	p-value	Interpretation
Strategy	2.635	2	0.268	Zero restriction not rejected
Objective	3.418	3	0.332	Zero restriction not rejected
Risk	23.452	1	0.0000013	Zero restriction rejected
Return	0.919	1	0.338	Zero restriction not rejected
Information	0.057	1	0.811	Zero restriction not rejected
Region of operation	4.024	2	0.134	Zero restriction not rejected

The fit statistics (reported in Appendix Table 12) reinforce the same point. For most attributes, the AIC and BIC values of the zero-constrained models are similar to, and in some cases slightly better than, those of the baseline models. For risk, however, both AIC and BIC worsen notably when the risk coefficient is fixed to zero. This implies that risk remains an important determinant of choice behaviour for both classes and should not be interpreted as an ignored attribute.

5.1.4 Class membership and covariates

Class membership probabilities were modelled as a function of the stated ANA indicator and a set of behavioural and demographic covariates. The full class membership estimates are reported in Appendix Table 16.

A clear pattern emerges across the specifications. The variable capturing whether the proposed funds feel better than respondents' current fund investments is the only covariate that is consistently statistically significant. Its coefficient is negative across the models, indicating that respondents who evaluate the proposed funds more favourably relative to their current investments are less likely to belong to Class B. By contrast, the stated ANA indicator, anticipated regret, gender, deliberation, and trust do not show robust significance across models. This suggests that self-reported non-attendance and most behavioural covariates do not strongly predict membership in the inferred non-attendance class. Overall, these results suggest that class membership is not primarily explained by demographic characteristics or stated ANA alone. Instead, the most consistent predictor is respondents' evaluation of the proposed investment alternatives relative to their current investments.

5.1.5 Stated versus inferred ANA

To provide a more intuitive assessment of the relationship between stated and inferred attribute non-attendance, Table 6 reports the average posterior probability of belonging to the ANA class (Class B) for individuals reporting and not reporting non-attendance, along with the corresponding standard deviations. The final column reports the difference in mean probabilities between the two groups.

Table 6. Average probability of Class B by stated ANA (investor dataset)

Attribute	ANA = 0 (Mean)	ANA = 0 (SD)	ANA = 1 (Mean)	ANA = 1 (SD)	Difference
Strategy	0.134	0.302	0.107	0.280	-0.027
Objective	0.115	0.282	0.156	0.334	0.041
Risk	0.290	0.384	0.343	0.376	0.053
Return	0.135	0.307	0.069	0.218	-0.066
Information	0.110	0.278	0.174	0.346	0.064
Region of operation	0.128	0.300	0.106	0.262	-0.022

In the latent class model, Class B represents individuals exhibiting attribute non-attendance, while Class A corresponds to full attendance. If stated ANA were a reliable indicator of behaviour, individuals reporting non-attendance would be expected to exhibit higher probabilities of belonging to Class B. However, the results reveal no consistent relationship across attributes. For example, for risk, individuals reporting non-attendance exhibit a higher average probability of belonging to Class B (0.343 compared with 0.290), corresponding to a positive difference of 0.053, indicating partial alignment between stated and inferred behaviour. A similar pattern is observed for information. In contrast, for return, the relationship is reversed, with individuals reporting non-attendance exhibiting a substantially lower probability of belonging to Class B (0.069 compared with 0.135), resulting in a negative difference of -0.066. Smaller negative differences are also observed for strategy and region of operation. The relatively large standard deviations indicate substantial heterogeneity within both groups. Overall, these findings indicate that the relationship between stated and inferred ANA is **weak and attribute-specific**, and that self-reported non-attendance does not consistently reflect underlying decision behaviour.

This pattern is further illustrated in Figure 5, which shows the distribution of posterior probabilities of belonging to Class B across attributes. The distributions exhibit substantial overlap between individuals reporting and not reporting non-attendance, reinforcing the lack of a clear relationship between stated and inferred ANA.

To further assess this relationship, posterior probabilities of belonging to Class B are correlated with the corresponding stated ANA indicators. The results show limited alignment between the two measures. For most attributes, the correlations are small and statistically insignificant. The only attribute with a statistically

significant correlation is information, with a small positive coefficient ($\rho = 0.077$, $p = 0.029$). The return attribute shows a weak negative relationship that is only borderline significant ($\rho = -0.068$, $p = 0.053$). For strategy, objective, risk, and region of operation, the estimated relationships are close to zero and clearly insignificant.

Table 7 reports these correlation results.

Table 7. Correlation between stated ANA and Class B probability (investor dataset)

Attribute	Correlation	p-value
Strategy	-0.030	0.396
Objective	0.055	0.118
Risk	0.025	0.472
Return	-0.068	0.053
Information	0.077	0.029
Region of operation	-0.028	0.434

Taken together, the results indicate that stated and inferred ANA do not strongly coincide in the investor dataset. This supports the view that self-reported simplification and behaviour inferred from choices capture related but distinct dimensions of decision-making. Additional regression analyses examining the relationship between the treatment and attribute non-attendance do not indicate statistically significant effects (see Appendix Table 15).

5.1.6 Summary of main findings

The descriptive and model-based results provide a consistent picture of attribute processing behaviour in the investor dataset. The following conclusions emerge. First, stated ANA is present across all attributes, but remains moderate in magnitude, with most respondents ignoring at most one attribute. Second, the latent class models consistently identify a small secondary class, indicating meaningful heterogeneity in decision strategies. Third, the comparison between baseline and zero-constrained models suggests that inferred ANA is plausible for several attributes, but not for risk, which remains consistently attended across specifications. Fourth, among the class membership covariates, only respondents' evaluation of the proposed funds relative to their current investments shows a robust and consistent relationship with class allocation. Finally, the association between stated and inferred ANA is weak overall, suggesting that self-reported and behaviourally inferred non-attendance should not be treated as equivalent.

5.1.7 Robustness check: restricted stated ANA sample

As a robustness check, the stated ANA variables were reconstructed by retaining only respondents who selected one of the first four substantive reasons for ignoring an attribute. These reasons include that the attribute was not important, made the choice easier, was perceived as unrealistic, or should not be weighed against the other attributes. Respondents selecting "don't know" were excluded

from this version of the analysis as a robustness exercise to examine the sensitivity of the results to the construction of stated ANA. As a result, the effective sample size decreases from 804 to 730 respondents.

The descriptive pattern remains very similar to the main analysis. The share of respondents reporting ANA decreases only slightly across attributes, by around 1 to 2 percentage points. The average number of ignored attributes also falls modestly, from 1.17 in the main specification to 1.09 in the restricted version. This indicates that the overall pattern of stated ANA is not driven by respondents who were uncertain about their reasons. The model-based results are also broadly consistent with the main findings. The zero restriction for risk is again rejected, confirming that risk should not be interpreted as an ignored attribute. In addition, the restricted specification yields some evidence of inferred ANA for region of operation, suggesting that this result is somewhat sensitive to how stated ANA is constructed (see Appendix Table 13 for detailed model fit statistics).

Overall, this robustness check confirms that the main conclusions remain unchanged: attribute non-attendance is present but selective, and risk remains a consistently attended attribute.

5.1.8 Extension: Multi-class ($n+1$) specification

As an exploratory extension beyond the main and robustness analyses, an $n + 1$ latent class model was estimated for the investor dataset. In this specification, one class represents full attribute attendance, while the remaining classes capture non-attendance of individual attributes, including strategy, objective, risk, return, information, and region of operation.

The estimated class probabilities indicate heterogeneity across these patterns (see Appendix Table 14). The largest class corresponds to non-attendance of information (0.3442). This is followed by non-attendance of risk (0.1928) and objective (0.1764), while the remaining classes are comparatively small, including region of operation (0.0777), return (0.0313), and strategy (0.0066). The relatively large share associated with risk non-attendance should be interpreted with caution, as the two-class results provide strong evidence that risk is not ignored.

In terms of model fit, the multi-class specification achieves a log-likelihood of -3730.49, with AIC = 7606.99 and BIC = 8080.04, which are numerically improved relative to the two-class models. However, these improvements must be interpreted cautiously due to estimation issues. The model converged with singular convergence, no covariance matrix could be computed, and several coefficients reached implausibly large magnitudes. In addition, one component yielded a likelihood of negative infinity for at least one individual.

These issues indicate that the multi-class specification is overfitting the data rather than identifying stable behavioural segments. While the model allows for a richer representation of attribute non-attendance, the lack of estimation stability prevents reliable interpretation. For this reason, the $n + 1$ specification is treated as an

exploratory extension, and the two-class models remain the main basis for interpretation.

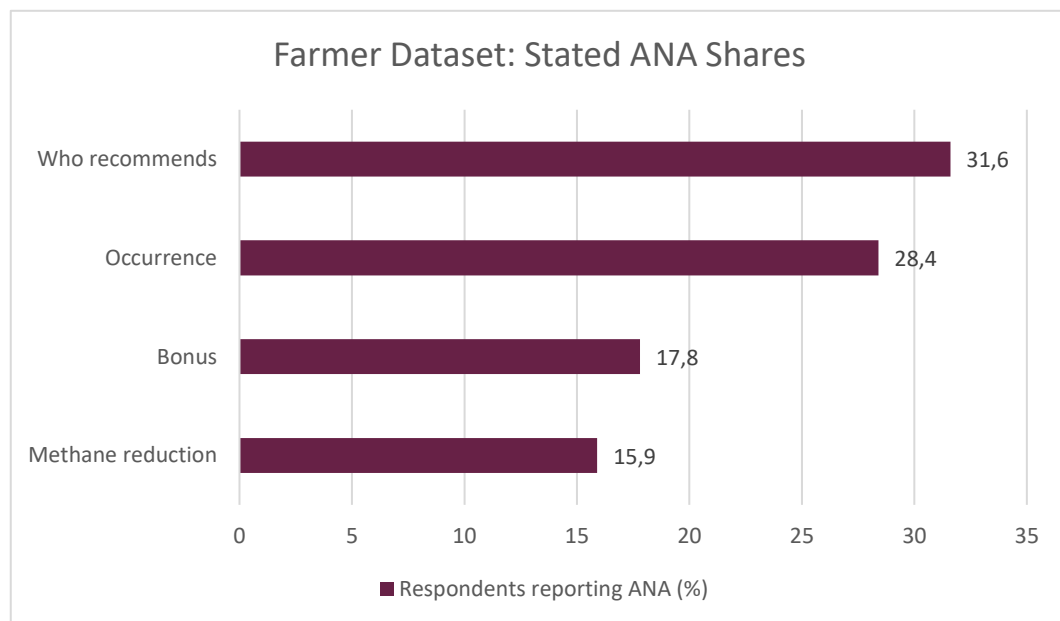
5.2 Farmer data

5.2.1 Stated attribute non-attendance

We begin by examining stated attribute non-attendance (ANA) based on farmers' self-reported behaviour.

The proportion of respondents reporting non-attendance varies across attributes. The highest shares are observed for the “Who recommends” attribute (31.6%) and Occurrence (28.4%), while lower levels are reported for methane reduction (15.9%) and bonus payments (17.8%). These results indicate that farmers do not ignore all attributes equally, but instead selectively simplify the decision task. Figure 2 shows the share of respondents reporting attribute non-attendance for each attribute in the farmer dataset.

Figure 2. Stated ANA shares across attributes in the farmer dataset.



Note: Shares are reported separately for each attribute. Respondents could report non-attendance for multiple attributes; therefore, percentages do not sum to 100%.

At the respondent level, the average number of ignored attributes is approximately **0.94**, indicating that simplification is present but relatively limited. Most respondents ignore none or only one attribute, while ignoring multiple attributes is less common. The relatively higher levels of non-attendance for recommendation-related attributes suggest that advisory aspects may be perceived as less central to the decision compared to economic and environmental outcomes.

Overall, these results suggest that stated ANA is present in the farmers dataset, but remains moderate in magnitude and varies across attributes.

5.2.2 Baseline latent class results

To account for preference heterogeneity and potential attribute non-attendance (ANA), two-class latent class models were estimated separately for each attribute. Across all baseline specifications, the class allocation results are highly stable. The dominant class (Class A) consistently represents approximately 78–79% of respondents, while the secondary class (Class B) accounts for around 21–22%.

Table 8 reports the class allocation shares for each attribute.

Table 8. Estimated class shares (baseline latent class models, farmer dataset)

Attribute	Class A	Class B
Who recommends	78.5%	21.5%
Occurrence	78.5%	21.5%
Methane reduction	78.8%	21.2%
Bonus	78.4%	21.6%

The reported values represent the estimated average class membership probabilities (class shares) from each attribute-specific baseline model.

These results indicate the presence of a smaller group of respondents who follow a different decision rule from the majority. This pattern suggests meaningful heterogeneity in decision-making behaviour, but does not by itself provide evidence of attribute non-attendance, which is formally assessed using zero-constrained models.

5.2.3 Baseline versus zero-constrained models

To test for inferred attribute non-attendance (ANA), each baseline model is compared with a corresponding zero-constrained specification, in which the coefficient of the relevant attribute is fixed to zero in Class B.

Table 9 reports the likelihood ratio (LR) test results for these comparisons.

Table 9. Likelihood ratio test results (farmer dataset)

Attribute	LR statistic	df	p-value	Interpretation
Who recommends	7.18	3	0.066	Not rejected
Occurrence	1.61	1	0.204	Not rejected
Methane reduction	5.76	3	0.124	Not rejected
Bonus	17.39	1	0.00003	Rejected

The results show a consistent pattern across attributes. For most attributes, the zero restriction is not rejected, indicating limited support for inferred ANA. A marginal effect is observed for *who recommends* ($p = 0.066$), but this remains statistically insignificant at conventional levels. In contrast, for the *bonus* attribute, the zero restriction leads to a clear and statistically significant deterioration in model fit. This indicates that bonus remains an important determinant of choice behaviour and should not be interpreted as an ignored attribute. Overall, these findings suggest that inferred ANA is limited in the farmers dataset and that most attributes remain relevant for decision-making.

5.2.4 Class membership and covariates

Class membership probabilities in the farmer models are specified as a function of the stated ANA indicator and a set of demographic and behavioural covariates, including gender, age, risk-taking behaviour, and generalised trust. Across the specifications, the included covariates do not show a strong or consistent relationship with class allocation. The stated ANA indicator varies across models and is only statistically significant in the bonus specification, where it is positively associated with Class B membership. In contrast, gender shows a consistently positive effect and is statistically significant in several models, indicating that female respondents are more likely to belong to Class B.

Age, risk-taking behaviour, and generalised trust do not appear to meaningfully explain class membership. Overall, these results suggest that observable characteristics only weakly explain variation in class allocation, which may instead reflect unobserved heterogeneity in how farmers process the choice tasks (see Appendix Table 17 for detailed estimates).

5.2.5 Stated versus inferred ANA

To assess the relationship between stated and inferred attribute non-attendance (ANA), posterior probabilities of belonging to the ANA class (Class B) were compared between individuals who reported ignoring an attribute (ANA = 1) and those who did not (ANA = 0).

Table 10 presents the average posterior probabilities by stated ANA status.

Table 10. Average probability of Class B by stated ANA (farmer dataset)

Attribute	ANA = 0 (Mean)	ANA = 0 (SD)	ANA = 1 (Mean)	ANA = 1 (SD)	Difference
Who recommends	0.226	0.063	0.191	0.061	-0.035
Occurrence	0.222	0.066	0.196	0.049	-0.025
Methane reduction	0.194	0.060	0.306	0.059	+0.112
Bonus	0.192	0.062	0.325	0.074	+0.133

The relatively small standard deviations indicate less dispersion in posterior probabilities compared with the investor dataset, although clear differences across attributes remain. Overall, the distributions suggest that the relationship between stated and inferred ANA varies substantially depending on the attribute being considered. The results reveal substantial variation across attributes. For bonus and methane reduction, individuals reporting non-attendance exhibit higher probabilities of belonging to Class B, indicating a degree of alignment between stated and inferred behaviour. This pattern is also illustrated in Figure 6, which shows the distribution of posterior class probabilities across attributes in the farmer dataset. In contrast, for advisor (who recommends) and occurrence, the relationship is weak and even negative, suggesting that self-reported non-attendance does not consistently correspond to behaviour inferred from the choice data. This pattern is further supported by correlation analysis.

Table 11 reports the correlation between stated ANA and posterior class probabilities.

Table 11. Correlation between stated ANA and Class B probability (farmer dataset)

Attribute	Correlation	p-value
Who recommends	-0.250	<0.001
Occurrence	-0.183	<0.01
Methane reduction	0.566	<0.001
Bonus	0.623	<0.001

Taken together, the results suggest that stated and inferred ANA align more strongly for economically and environmentally central attributes such as bonus and methane reduction, while little alignment is observed for advisory-related attributes. This reinforces the view that stated ANA does not uniformly correspond to behaviour inferred from observed choices.

This pattern is consistent with the model comparison results reported in Appendix Table 18. The zero-constrained models do not improve model fit for most attributes, and perform substantially worse for the bonus attribute, suggesting that respondents do not generally ignore attributes in the manner implied by stated ANA. Overall, these findings indicate that stated and inferred ANA capture related but distinct dimensions of decision-making. While the two measures are aligned for some attributes, they diverge for others, suggesting that stated ANA may reflect simplification strategies or ex-post rationalisation rather than actual decision processes.

5.2.6 Summary of main findings

The results for the farmers dataset provide a consistent picture of how respondents process attributes in the choice experiment. Stated attribute non-attendance (ANA) is present across all attributes, but remains moderate in magnitude, with farmers typically ignoring fewer than one attribute on average. At the same time, the latent class models reveal a stable minority group of respondents who appear to follow a different decision rule, highlighting the presence of heterogeneity in preferences and decision-making behaviour.

In terms of inferred ANA, the evidence is limited overall, with clear support emerging primarily for the bonus attribute. This suggests that respondents may still partially consider most attributes during decision-making, even when they report ignoring them. The comparison between stated and inferred ANA further shows that the relationship between the two measures is generally weak, but becomes more pronounced for specific attributes such as bonus and methane reduction. Taken together, these findings indicate that attribute non-attendance in the farmers dataset is selective and attribute-specific, rather than reflecting a general simplification strategy applied uniformly across respondents.

5.2.7 Robustness check: restricted stated ANA sample

As a robustness check, the farmer analysis was repeated using a restricted definition of stated ANA. Only respondents who selected one of the first four substantive reasons for ignoring an attribute were retained: the attribute was not important, made the choice easier, had unrealistic levels, or should not be weighted against other attributes. Respondents selecting “don’t know” or “prefer not to say” were excluded from this version of the analysis as a robustness exercise to examine the sensitivity of the results to the construction of stated ANA. As a result, the effective sample size decreases from 320 to 279 respondents.

The descriptive pattern remains broadly similar to the main analysis, although stated non-attendance decreases slightly across all attributes. The average number of ignored attributes falls from approximately 0.94 to 0.84, suggesting that excluded respondents were somewhat more likely to report non-attendance. The model-based results remain stable (see Appendix Table 19). Class allocation shares are largely unchanged, and the likelihood ratio tests continue to show strong evidence against the zero restriction for the bonus attribute. This confirms that bonus remains an important determinant of farmers’ choices rather than an ignored attribute. The relationship between stated and inferred ANA is also consistent with the main results: most attributes show weak or inconsistent alignment, while bonus displays a particularly strong association, with a correlation of approximately 0.70.

Overall, the robustness check indicates that the main findings are not driven by uncertain or low-quality stated ANA responses. It reinforces the conclusion that attribute non-attendance in the farmers dataset is selective and attribute-specific rather than a general behavioural pattern.

6. Discussion and Conclusion

The results from both the investor and farmer datasets provide consistent evidence that attribute non-attendance (ANA) is present, but not dominant. In both contexts, respondents do not ignore the choice tasks entirely; instead, they simplify decisions by focusing on a subset of attributes. Descriptive findings show that the share of respondents reporting non-attendance remains moderate across attributes, and the average number of ignored attributes is close to one. This suggests that simplification is common, but not extreme, and that most respondents continue to engage meaningfully with the choice tasks. This finding is consistent with previous research suggesting that respondents often simplify complex choice tasks without completely ignoring the available information (Gonçalves et al., 2022; Hensher, 2006). Rather than reflecting random decision-making, attribute non-attendance appears to represent a selective simplification strategy.

At the same time, ANA is clearly not uniform across attributes. In the investor dataset, most attributes can be constrained to zero without significantly worsening model fit, indicating that non-attendance is plausible for these dimensions. However, risk stands out as a clear exception, as imposing the zero restriction leads to a notable deterioration in fit. This highlights that risk remains a central determinant of investment decisions and is consistently attended to by respondents. A similar pattern emerges in the farmer dataset. While most attributes do not provide strong evidence against the zero restriction, the bonus attribute does. Likelihood ratio tests show that constraining the bonus coefficient to zero significantly worsens model fit, suggesting that financial incentives are a key driver of farmers' choices. Taken together, these findings indicate that respondents simplify decisions selectively, while continuing to attend to attributes that are most relevant within the given context. This supports the argument that ANA is not uniform across attributes, as previously documented by Hensher et al. (2005) and Scarpa et al. (2009). Attributes perceived as particularly important are less likely to be ignored, even when respondents simplify other aspects of the choice task.

The latent class results further underline the presence of preference heterogeneity in both datasets. In each case, a dominant class accounts for the majority of respondents, while a smaller secondary class captures individuals who follow a different decision pattern. However, the comparison between baseline and zero-constrained models suggests that this secondary class cannot always be interpreted purely as an ANA class. Instead, it appears to reflect broader differences in decision-making behaviour. This implies that heterogeneity in preferences and simplification strategies coexist, but are not perfectly aligned. This interpretation is consistent with Hensher and Greene (2010) who argue that latent classes may capture both attribute-processing behaviour and broader preference heterogeneity. As a result, inferred ANA cannot always be separated cleanly from differences in underlying preferences.

A central objective of this study was to compare stated and inferred measures of ANA. The results consistently show that the relationship between these two measures is weak and highly attribute-specific. In the investor dataset, correlations between stated non-attendance and posterior class probabilities are small and largely insignificant, and differences in mean probabilities across groups do not reveal a consistent pattern. In the farmer dataset, some attributes show stronger alignment, particularly methane reduction and bonus, where respondents reporting non-attendance exhibit higher probabilities of belonging to the secondary class. However, for other attributes, the relationship remains weak or even negative. Overall, these findings indicate that self-reported non-attendance is only weakly associated with behaviour inferred from observed choices. Stated and inferred measures appear to capture related but distinct aspects of decision-making, and should therefore not be treated as interchangeable. This finding is broadly consistent with previous studies reporting limited correspondence between stated and inferred ANA measures (Campbell et al., 2011; Hensher & Greene, 2010). One explanation is that respondents may not be fully aware of the simplification strategies they employ, or that self-reported non-attendance reflects perceptions of behaviour rather than actual decision processes.

The role of observed individual characteristics in explaining class membership appears limited. In the investor dataset, the only covariate that is consistently statistically significant is whether respondents perceived the proposed funds as better than their current fund investments. In contrast, other variables, such as gender, deliberation, and trust, do not exhibit stable effects. In the farmer dataset, gender shows some association with class membership, while age, risk-taking behaviour, and generalised trust do not show strong or consistent relationships. This suggests that ANA and decision strategies are not primarily driven by observable demographic or behavioural characteristics, but may instead reflect more subtle aspects of how individuals process information in complex choice environments. The generally weak explanatory power of observable characteristics is also consistent with the view that attribute-processing behaviour is influenced by unobserved cognitive mechanisms that are difficult to capture through standard demographic variables alone.

Robustness analyses conducted for both datasets further support the stability of the findings. In the investor dataset, reconstructing the stated ANA variables using only respondents who selected substantive reasons for ignoring attributes slightly reduced the sample size and lowered the overall level of reported ANA. However, the descriptive and model-based results remained broadly unchanged. Risk continued to be consistently attended, while most other attributes could plausibly be ignored without substantially worsening model fit. Some additional evidence of inferred ANA emerged for region of operation, suggesting limited sensitivity to how stated ANA is defined. Similarly, in the farmer dataset, excluding respondents who reported less meaningful reasons for ignoring attributes reduced the sample size and slightly lowered the average number of ignored attributes, but the main findings remained stable. Class allocation shares changed little, likelihood ratio tests continued to reject the zero restriction only for the bonus

attribute, and the relationship between stated and inferred ANA remained largely consistent. Overall, these robustness checks strengthen confidence that the main conclusions are not driven by low-quality responses or by the specific construction of the stated ANA measures.

A comparison of the two datasets highlights both similarities and important differences. In both contexts, ANA is selective rather than universal, and the majority of respondents attend to most attributes. However, the attributes that remain consistently attended differ across settings. In investment decisions, risk plays a central role and is unlikely to be ignored, whereas in the agricultural context, monetary incentives in the form of bonus payments appear to be the key driver of choices. This suggests that attribute importance is context-dependent and closely linked to the nature of the decision problem. This reinforces findings from previous ANA research that attribute processing is context dependent and varies according to the decision environment and the salience of specific attributes (Caputo et al., 2018; Yangui et al., 2025).

From a policy perspective, these findings have important implications for the design of choice-based interventions and surveys. The results suggest that individuals do not process all available information equally, but instead focus on a subset of attributes that are most relevant to them. This implies that policies relying on complex information provision may be less effective if key attributes are overlooked. For example, in the investment context, risk information appears to be consistently attended and should therefore be clearly communicated, while less critical attributes may require simplification or prioritisation. Similarly, in the agricultural context, financial incentives such as bonus payments play a central role in decision-making, indicating that policy instruments targeting economic outcomes are likely to be more effective than those relying solely on informational or advisory components. More broadly, the weak alignment between stated and inferred non-attendance suggests that relying solely on self-reported measures may lead to misleading conclusions about behaviour. Policymakers and researchers should therefore complement survey-based approaches with behavioural evidence when designing and evaluating interventions.

These findings should be interpreted in light of several limitations. The analysis relies on a two-class latent class specification, which provides a simplified representation of heterogeneity and may not capture more complex patterns of behaviour. In addition, stated ANA is based on self-reported responses, which may be subject to reporting bias. The inferred measure of non-attendance is model-dependent and relies on the imposed zero restrictions within the latent class framework. Finally, the results are based on two specific datasets, and their generalisability to other contexts may be limited. In addition, both datasets include split-sample treatments that were pooled in the latent class analysis. While pooling increases statistical power, future research could explicitly test for scale heterogeneity across treatment groups before estimating pooled latent class models to determine whether observed differences reflect preference

heterogeneity, attribute-processing behaviour, or differences in decision consistency across treatment conditions.

Overall, the findings support the growing literature suggesting that attribute non-attendance should be understood as a selective and context-dependent behaviour rather than a universal decision rule. Respondents simplify complex choice tasks by selectively ignoring some attributes while continuing to focus on those that are most salient within the decision context. At the same time, the weak relationship between stated and inferred measures indicates that self-reported behaviour does not fully capture underlying decision processes. Taken together, these results highlight the importance of combining descriptive and model-based approaches when analysing attribute processing in discrete choice experiments, as relying on a single measure of non-attendance may provide an incomplete picture of behaviour.

References

- Alemu, M. H., Mørkbak, M. R., Olsen, S. B., & Jensen, C. L. (2013). Attending to the Reasons for Attribute Non-attendance in Choice Experiments. *Environmental and Resource Economics*, 54(3), 333–359. <https://doi.org/10.1007/s10640-012-9597-8>
- Alho, E. (2017). Attention to risk and return: Choice experiment of the stated and inferred use of investment attributes. *Applied Economics and Finance*, 4(1), 43–52.
- Alpizar, F., Carlsson, F., & Martinsson, P. (2001). *Using choice experiments for non-market valuation*. <https://idl-bnc-idrc.dspacedirect.org/bitstreams/e255ca7e-604f-46ef-8cb6-0c4ed3241b89/download>
- Campbell, D., Hensher, D. A., & Scarpa, R. (2011). Non-attendance to attributes in environmental choice analysis: A latent class specification. *Journal of Environmental Planning and Management*, 54(8), 1061–1076. <https://doi.org/10.1080/09640568.2010.549367>
- Caputo, V., Van Loo, E. J., Scarpa, R., Nayga, R. M., & Verbeke, W. (2018). Comparing Serial, and Choice Task Stated and Inferred Attribute Non-Attendance Methods in Food Choice Experiments. *Journal of Agricultural Economics*, 69(1), 35–57. <https://doi.org/10.1111/1477-9552.12246>
- Carlsson, F., Kataria, M., & Lampi, E. (2010). Dealing with Ignored Attributes in Choice Experiments on Valuation of Sweden's Environmental Quality Objectives. *Environmental and Resource Economics*, 47(1), 65–89. <https://doi.org/10.1007/s10640-010-9365-6>
- Edenbrandt, A. K., Lagerkvist, C.-J., Lüken, M., & Orquin, J. L. (2022). Seen but not considered? Awareness and consideration in choice analysis. *Journal of Choice Modelling*, 45, 100375.
- Gonçalves, T., Lourenço-Gomes, L., & Pinto, L. M. C. (2022). The role of attribute non-attendance on consumer decision-making: Theoretical insights and empirical

- evidence. *Economic Analysis and Policy*, 76, 788–805.
<https://doi.org/10.1016/j.eap.2022.09.017>
- Gottlieb, U., & Rommel, J. (2026). The role of advisory provision and policy objective information in dairy farmers' preferences for methane-reducing feed additives. *European Review of Agricultural Economics*, 53(1), 446–472.
- Greene, W. H., & Hensher, D. A. (2003). A latent class model for discrete choice analysis: Contrasts with mixed logit. *Transportation Research Part B: Methodological*, 37(8), 681–698.
- Hensher, D. A. (2006). How do respondents process stated choice experiments? Attribute consideration under varying information load. *Journal of Applied Econometrics*, 21(6), 861–878. <https://doi.org/10.1002/jae.877>
- Hensher, D. A., & Greene, W. H. (2010). Non-attendance and dual processing of common-metric attributes in choice analysis: A latent class specification. *Empirical Economics*, 39(2), 413–426. <https://doi.org/10.1007/s00181-009-0310-x>
- Hensher, D. A., Rose, J., & Greene, W. H. (2005). The implications on willingness to pay of respondents ignoring specific attributes. *Transportation*, 32(3), 203–222.
<https://doi.org/10.1007/s11116-004-7613-8>
- Hensher, D. A., Rose, J. M., & Greene, W. H. (2015). *Applied choice analysis*. Cambridge university press.
- Hoyos, D. (2010). The state of the art of environmental valuation with discrete choice experiments. *Ecological Economics*, 69(8), 1595–1603.
- McFadden, D. (1974). *Analysis of qualitative choice behavior*. Zarembka, P.(ed.): *Frontiers in Econometrics*. Academic Press. New York, NY.
- Rakotonarivo, O. S., Schaafsma, M., & Hockley, N. (2016). A systematic review of the reliability and validity of discrete choice experiments in valuing non-market environmental goods. *Journal of Environmental Management*, 183, 98–109.

- Scarpa, R., Gilbride, T. J., Campbell, D., & Hensher, D. A. (2009). Modelling attribute non-attendance in choice experiments for rural landscape valuation. *European Review of Agricultural Economics*, 36(2), 151–174.
- Schulze, C., Zagórska, K., Häfner, K., Markiewicz, O., Czajkowski, M., & Matzdorf, B. (2024). Using farmers' ex ante preferences to design agri-environmental contracts: A systematic review. *Journal of Agricultural Economics*, 75(1), 44–83. <https://doi.org/10.1111/1477-9552.12570>
- Simon, H. A. (1955). A behavioral model of rational choice. *The Quarterly Journal of Economics*, 99–118.
- Train, K. E. (2009). *Discrete choice methods with simulation*. Cambridge university press. <https://books.google.com/books?hl=en&lr=&id=4yHaAgAAQBAJ&oi=fnd&pg=PR7&dq=Discrete+Choice+Methods+with+Simulation+Train+2009&ots=eG1eHffYdx&sig=HLYGEHgxSzJSbGryiNC8aESToc8>
- Villanueva, A. J., Perni, A., & Barreiro-Hurle, J. (2026). Choice experiments on land managers' participation in environmental programs: A systematic review and meta-analysis of estimate validity. *American Journal of Agricultural Economics*, ajae.70073. <https://doi.org/10.1002/ajae.70073>
- Yangui, A., Akaichi, F., & Gil, J. M. (2025). Investigating the Effect of Attribute Non-Attendance in Different Elicitation Formats: Single Discrete Choice, Rank-Order Discrete Choice, and Best Worst Scaling. *Agricultural Economics*, 56(6), 1079–1102. <https://doi.org/10.1111/agec.70053>

Popular science summary

People often have to make choices involving multiple pieces of information. For example, farmers may choose between different agricultural practices, while investors may decide which financial products to invest in. These decisions are often studied using surveys where individuals are asked to choose between alternatives described by several characteristics.

A common assumption in such studies is that people carefully consider all the information provided. However, in reality, individuals may simplify complex decisions by ignoring some of the information. This behaviour is known as attribute non-attendance.

This thesis studies how often people ignore information when making decisions and whether this behaviour differs across contexts. Two different datasets are analysed: one based on farmers making decisions about agricultural practices, and another based on investors choosing between sustainable investment options.

The analysis shows that people do not always consider all available information. Instead, they focus on certain key aspects. For farmers, financial incentives appear to play an important role, while for investors, risk is a key factor influencing decisions.

The study also compares what people say about their decision-making with what can be observed from their choices. The results suggest that people's self-reports do not always fully reflect their actual behaviour.

Overall, the findings highlight that simplifying decision strategies are common and should be taken into account when designing policies and interpreting survey-based studies. Understanding how people process information can help improve both policy design and the reliability of economic analysis.

Appendix

Figure 3. Example of a choice task from the farmer dataset

	Feed additive A	Feed additive B	None of these. Keep current practices.
Who gives advice	Experienced farmer	Veterinarian	
How often is advice given	Two times per year	Two times per year	
Reduction of methane emissions	30% reduction	10% reduction	
Bonus per kg of milk	0 öre	5 öre	

Figure 4. Example of a choice task from the investor dataset

Consider carefully all information in the choice situation. Choose if and how you would invest 20% of your monthly investments into equity funds if you were to choose in real life.

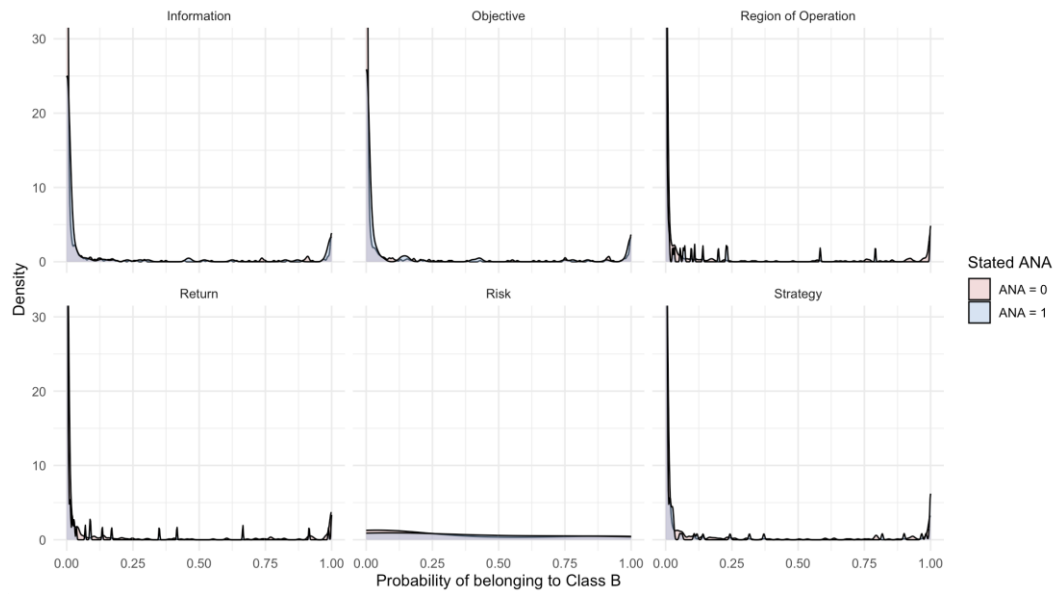
	Fund A	Fund B	Neither A nor B.
<i>Environmental objective</i>	 Transition to a circular economy	 Pollution prevention and control	
<i>Management strategy</i>	 Exclude	 Select	
<i>Region of operations</i>	 Global	 Sweden	
<i>Environmental performance information</i>	 Environmental score by a third party	 Assured by an auditor	
<i>Risk classification</i>	 second highest risk class 6	 upper medium risk class 5	
<i>Expected annual return</i>	 Return 9%	 Return 6%	

I choose fund A

I choose fund B

If these are the only funds to choose from, I choose none

Figure 5. Distribution of ANA Class Probabilities by Attribute (Investor data)



Note: The y-axis is truncated to improve visual readability because posterior probabilities are highly concentrated near zero.

Figure 6. Distribution of ANA Class Probabilities by Attribute (Farmer data)

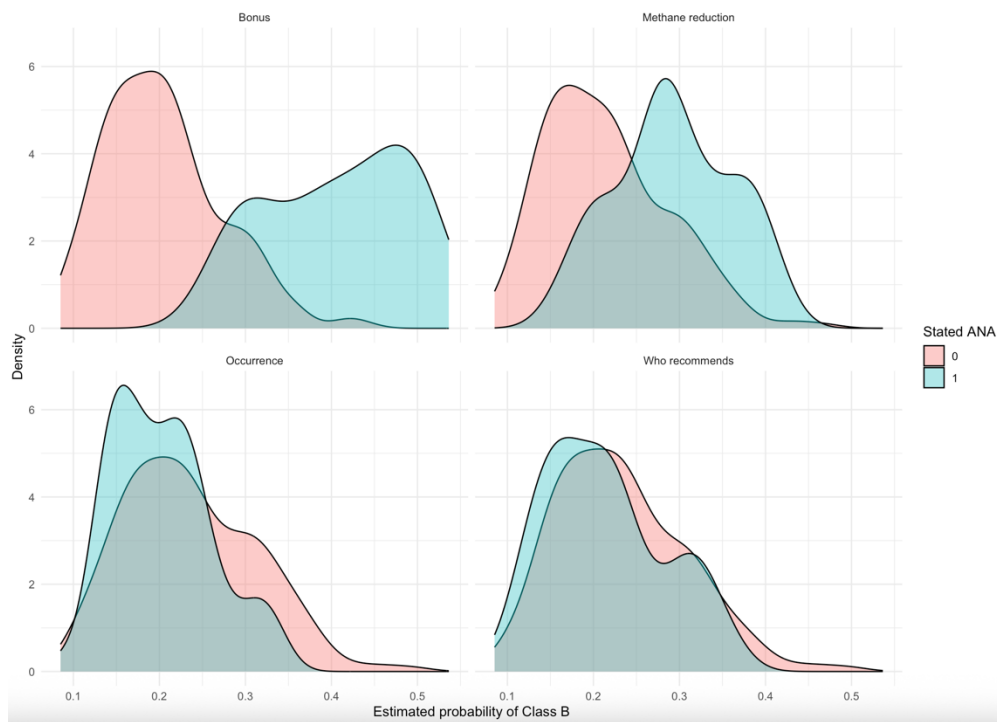


Table 12. Model fit statistics (two-class latent class models, investor dataset)

Attribute	Model	LL	AIC	BIC
Strategy	Baseline	-3947.806	7953.612	8141.536
Strategy	Zero	-3949.123	7952.247	8127.210
Objective	Baseline	-3948.032	7954.063	8141.987
Objective	Zero	-3949.741	7951.481	8119.964
Risk	Baseline	-3953.703	7965.405	8153.328
Risk	Zero	-3965.429	7986.857	8168.300
Return	Baseline	-3946.319	7950.638	8138.561
Return	Zero	-3946.778	7949.556	8131.000
Information	Baseline	-3946.785	7951.570	8139.494
Information	Zero	-3946.814	7949.627	8131.070
Region of operation	Baseline	-3948.238	7954.476	8142.399
Region of operation	Zero	-3950.250	7954.500	8129.464

Table 13. Model fit statistics with restricted stated ANA (investor dataset, $N = 730$)

Attribute	Model	LL	AIC	BIC
Strategy	Baseline	-3573.115	7204.230	7389.390
Strategy	Zero	-3573.736	7201.471	7373.861
Objective	Baseline	-3573.549	7205.098	7390.257
Objective	Zero	-3574.887	7201.775	7367.780
Risk	Baseline	-3581.843	7221.687	7406.846
Risk	Zero	-3585.972	7227.944	7406.718
Return	Baseline	-3571.569	7201.138	7386.297
Return	Zero	-3572.139	7200.277	7379.052
Information	Baseline	-3573.048	7204.097	7389.256
Information	Zero	-3573.153	7202.307	7381.081
Region of operation	Baseline	-3572.137	7202.275	7387.434
Region of operation	Zero	-3575.290	7204.579	7376.969

Table 14. Estimated class probabilities ($n+1$ latent class model, investor dataset)

Class	Interpretation	Mean probability
Class 1	Full attendance	0.1709
Class 2	Ignore strategy	0.0066

Class 3	Ignore objective	0.1764
Class 4	Ignore risk	0.1928
Class 5	Ignore return	0.0313
Class 6	Ignore information	0.3442
Class 7	Ignore region of operation	0.0777

Table 15. Effect of T60 treatment on stated ANA (investor dataset)

Attribute	Coefficient (t60)	Odds Ratio	p-value
ANAAstr	0.048	1.049	0.777
ANAobj	-0.098	0.906	0.568
ANArisk	0.231	1.260	0.213
ANArete	0.135	1.144	0.452
ANAinf	-0.025	0.976	0.888
ANArege	-0.064	0.938	0.738

Notes: Coefficients are obtained from logistic regressions where each attribute-specific ANA indicator is regressed on the treatment indicator (T60). Odds ratios are reported as exponentiated coefficients. None of the estimated treatment effects are statistically significant at conventional levels.

Table 16. Class membership estimates (investor dataset)

Variable	Coefficient (Class B)	Significance
Stated ANA (attribute-specific)	≈ 0.00 to 0.10	ns
Anticipated regret (AR_f)	≈ -0.05 to -0.15	ns
Proposed funds feel better (Inv_Feel_Comp)	≈ -0.56 to -0.59	***
Female (gender_fem)	≈ -0.10 to -0.30	ns
Deliberation (del_f)	≈ 0.05 to 0.20	ns
Trust (TrustAll_f)	≈ 0.00 to 0.15	ns

Notes: *** $p < 0.01$, ** $p < 0.05$, * $p < 0.10$, ns = not statistically significant. Coefficients refer to the utility of belonging to Class B relative to Class A. Reported values summarise patterns across attribute-specific baseline models.

Table 17. Class membership estimates (farmer dataset)

Variable	Coefficient (Class B)	Significance
----------	-----------------------	--------------

Stated ANA (attribute-specific)	≈ -0.21 to 0.73	mixed (ns to **)
Female	≈ 0.65 to 0.72	* to **
Age	≈ 0.15 to 0.19	ns
Risk-taking	≈ 0.17 to 0.21	ns
Generalised trust	≈ 0.15 to 0.17	ns

Notes: *** $p < 0.01$, ** $p < 0.05$, * $p < 0.10$, ns = not statistically significant. Coefficients refer to the utility of belonging to Class B relative to Class A. Reported values summarise patterns across attribute-specific baseline models.

Table 18. Model fit statistics (two-class latent class models, farmer dataset)

Attribute	Model	LL	AIC	BIC
Advisor (Who recommends)	Baseline	-1582.278	3212.557	3302.996
Advisor (Who recommends)	Zero	-1585.869	3219.738	3310.178
Occurrence	Baseline	-1582.485	3212.969	3303.409
Occurrence	Zero	-1583.292	3214.584	3305.024
Methane reduction	Baseline	-1581.213	3210.426	3300.866
Methane reduction	Zero	-1584.093	3216.185	3306.625
Bonus	Baseline	-1580.451	3208.901	3299.341
Bonus	Zero	-1589.146	3226.291	3316.731

Table 19. Model fit statistics with restricted stated ANA (farmer dataset, $N = 279$)

Attribute	Model	LL	AIC	BIC
Advisor (Who recommends)	Baseline	-1365.150	2778.299	2865.448
Advisor (Who recommends)	Zero	-1366.813	2781.625	2868.774
Occurrence	Baseline	-1365.165	2778.330	2865.479
Occurrence	Zero	-1366.614	2781.227	2868.376
Methane reduction	Baseline	-1364.773	2777.545	2864.694
Methane reduction	Zero	-1367.774	2783.547	2870.697
Bonus	Baseline	-1362.727	2773.454	2860.603
Bonus	Zero	-1368.615	2785.230	2872.379

Publishing and archiving

Approved students' theses at SLU can be published online. As a student you own the copyright to your work and in such cases, you need to approve the publication. In connection with your approval of publication, SLU will process your personal data (name) to make the work searchable on the internet. You can revoke your consent at any time by contacting the library.

Even if you choose not to publish the work or if you revoke your approval, the thesis will be archived digitally according to archive legislation.

You will find links to SLU's publication agreement and SLU's processing of personal data and your rights on this page:

- <https://libanswers.slu.se/en/faq/228318>

☒ YES, I, Shahmeer Akhtar, have read and agree to the agreement for publication and the personal data processing that takes place in connection with this

☐ NO, I/we do not give my/our permission to publish the full text of this work. However, the work will be uploaded for archiving and the metadata and summary will be visible and searchable.