

Confirmation Bias Shapes Outlier Identification and Removal in Data Cleaning with Scatterplots

Shiyao Li , Akul Rana, Cindy Xiong Bearfield , and Emily Wall 

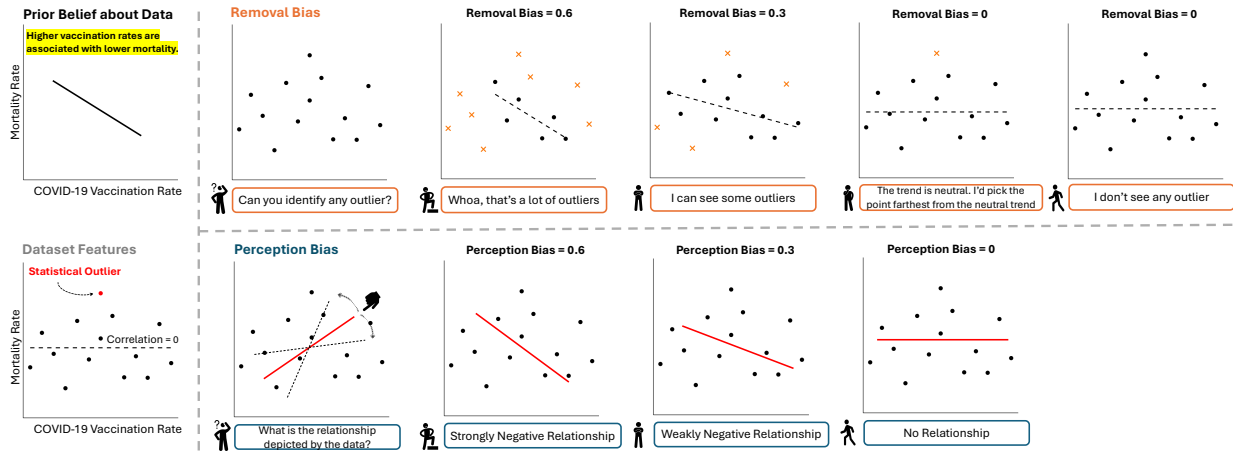


Figure 1: A schematic depiction of two manifestations of confirmation bias tested in this work: Removal Bias and Perception Bias

Abstract—Confirmation bias can shape how prior beliefs influence both the perception of scatterplot correlations and the identification of outliers during data cleaning. We characterize this bias in two stages: perception bias, in which people interpret a scatterplot's correlation in ways that align with their beliefs, and outlier bias, in which they selectively remove points to make the resulting data better match those beliefs. In a pre-registered crowdsourced study with 507 participants, we designed an interactive 'Find Outliers' task in which participants first expressed their prior beliefs about the relationship shown in a scatterplot before viewing it. Then, if they believed outliers were present, they identified and removed those outliers. Overall, 59% of participants removed data points. Among them, stronger prior beliefs predicted greater perception and outlier bias. They also tended to remove data points in ways that made the remaining data more consistent with their beliefs, especially when the outliers were less visually salient and the data were therefore more open to interpretation. Our findings show that data cleaning is not purely objective. Prior beliefs can guide how people visually identify and remove outliers. We therefore discuss how visualization systems can better support data cleaning workflows to mitigate confirmation bias.

Index Terms—Confirmation bias, data visualization, outlier removal

1 INTRODUCTION

Data cleaning directly affects data quality and, consequently, the validity of downstream results [26]. One key component of data cleaning involves identifying and removing outliers, which are data points that deviate significantly from other observations [25–27]. Such outliers can result from erroneous data entry, measurement error, or fraudulent responses. When left unaddressed, they can distort summary statistics, create misleading visual patterns, and lead to analytic results that fail to replicate. Making sound data cleaning decisions is, therefore, critical to rigorous data analysis.

However, the outlier removal process is not standardized across disciplines or analytic workflows. As a result, researchers face substantial degrees of freedom when deciding how to identify and treat outliers. For example, some may apply a standard deviation rule, removing points that fall more than 2 or 3 standard deviations from the mean; others may use percentile trimming, excluding the top and bottom one

or two percent of observations; others may cap extreme values at a threshold rather than remove them altogether. Manual inspection and intervention then become especially useful when statistical or automated methods can flag potential outliers but cannot determine how those points should be handled. In such cases, researchers may rely on domain expertise to exclude values they deem implausible, impossible, or inconsistent with known constraints. Prior data-cleaning and wrangling systems also reflect this human-in-the-loop workflow. For example, Potter's Wheel supports manual intervention through a spreadsheet-like interface where users can identify potential outliers, and iteratively apply transformations to correct, standardize, or remove those data points [45]. Wrangler allows analysts to select rows, columns, or cell text that may reflect outliers, then preview, edit, and apply suggested cleaning transformations such as reformatting and deleting data [29]; and Profiler combines statistical summaries with linked visualizations to highlight potential outliers so that analysts can inspect flagged values and decide how to handle them.

By making outliers visible, visualizations can help analysts detect obvious issues such as erroneous data entries before applying statistical procedures such as the standard deviation rule or percentile trimming [2, 22]. However, this initial visual assessment can also be problematic. Although visualizations can help analysts understand how to approach data cleaning, prior beliefs may also influence what people see and how they interpret the data [59], leading them to perceive patterns in ways that align with those beliefs and to make data cleaning decisions accordingly. This tendency is known in cognitive science

• Shiyao Li, Akul Rana, and Emily Wall are with Emory University. E-mails: {shiyao.li, akul.rana, emily.wall}@emory.edu.

• Cindy Xiong Bearfield is with Georgia Institute of Technology. E-mail: cxiong@gatech.edu.

Manuscript received xx xxx. 202x; accepted xx xxx. 202x. Date of Publication xx xxx. 202x; date of current version xx xxx. 202x. For information on obtaining reprints of this article, please send e-mail to: reprints@ieee.org. Digital Object Identifier: xx.xxxx/TVCG.202x.xxxxxxx

as confirmation bias: the tendency to seek, interpret, and even edit information in ways that support one's preexisting beliefs [42].

We posit that confirmation bias can impact data cleaning in two ways: perceptually and statistically. Perceptually, analysts may judge some data points as more or less likely to be outliers depending on whether those points make an apparent trend or correlation more consistent with their prior beliefs. Statistically, these biased perceptions may then carry over into outlier treatment decisions. Because data cleaning affords substantial degrees of freedom, analysts can choose whether and how to exclude outliers. For example, when applying a standard deviation rule, they may adopt a more lenient (e.g., 1.5 SD) or more stringent (e.g., 3 SD) threshold depending on whether excluding a particular point makes the resulting pattern more consistent with their belief or hypotheses. From there, confirmation bias can propagate through the data analysis pipeline. Here, visualizations, rather than serving as a helpful sanity check that informs statistical treatment, may instead bias subsequent decisions and reduce analytic rigor.

In this paper, we show that confirmation bias can influence data cleaning through both perceptual and statistical pathways, even in simple analytic settings. We designed a visual analytic task in which participants directly manipulated data and removed outliers when they deemed it appropriate, similar to the 'Find Anomalies' analytic task from Amar et al.'s 2005 framework [3]. We asked the research question:

How do prior beliefs and the visual characteristics of data shape people's interpretation of visualized data and their data cleaning decisions?

Our results show that participants were more likely to exclude data points when doing so made the resulting visualization more consistent with their prior beliefs. We further characterize the conditions under which people begin to rely on those beliefs when making outlier removal decisions. Together, these findings highlight a trade-off in using visualizations to inform analysis: although visualizations can support data inspection and improve analytic rigor, they can also shape perception in ways that allow confirmation bias to enter early in the analytic pipeline.

Because analysts often approach data with prior hypotheses, some degree of confirmation bias may be unavoidable. Although practices such as preregistration may help, they are not always feasible in fast-moving research contexts. Therefore, we hope the results of this paper can improve our understanding of confirmation bias, to inform future work on visualization system designs to preserve the benefits of visual inspection while helping reduce downstream bias.

2 RELATED WORKS

2.1 Statistical Outlier Identification Methods

Outlier identification is an essential step in quantitative analysis, yet the practices vary widely across different research domains. Drawing from Agunis et al. and Sullivan et al., we can characterize outlier identification from three dimensions: definitions, techniques, and handlings [2, 49].

Outliers can be defined in multiple ways – ranging from erroneous outliers due to data entry issues, to model-fit outliers that exert undue influence on statistical models. Aguinis et al., documented 14 mutually exclusive definitions, and 39 distinct identification techniques across a review of 232 organizational science research articles, highlighting the diversity and inconsistency in current practices. Their findings also underscore the importance of selecting outlier identification methods that are appropriate for the specific type of outlier and the underlying data structure. Common identification techniques for univariate data include identifying data points beyond 1.5 Interquartile range (IQR) or past center percentiles (e.g., 2.5th and 97.5th). For multivariate datasets, detection techniques often rely on distance-based metrics, such as identifying data points with studentized deleted residuals [7], squared Mahalanobis distances [40], or Cook's distances that exceed predefined thresholds [18]. Aguinis et al., also identified 20 handling strategies for these outliers, ranging from simple techniques such as correcting

erroneous values or removing extreme values, to more sophisticated transformation, such as Winsorizing, which replaces outliers with specific percentiles (e.g., the 5th or the 95th), or applying logarithmic transformations to reduce their influence on analysis [2].

Despite the diversity of outlier definitions, corresponding identification and handling methods, Sullivan et al., emphasized in their 2021 review that outlier handling is often inadequately reported in empirical research [49]. Across their review of 2,582 research articles published in top business journals, Winsorizing was the most frequently reported handling method, appearing in 24.7% of the articles. Additionally, 22.4% articles made no mention of outliers, while 21.6% mentioned identifying and handling outliers but failed to report the threshold used to define them. Given that even a small number of outliers can distort statistical significance—either by producing false positives or masking real relationships—the lack of rigor in current research practices underscores the need for greater consistency and transparency in the definition, identification, and handling of outliers. Additionally, prior research has shown that researchers are more likely to detect a statistically significant effect for a false hypothesis than to correctly find no effect in empirical studies [47]. This issue could be exacerbated by inconsistent outlier identification practices, which introduces ambiguity into empirical findings and may become a chronic issue in research, as researchers can cherry-pick thresholds that best support their desired hypotheses.

2.2 The Role of Visualization in Outlier Identification

Outlier identification commonly begins with data visualization, followed by quantitative methods such as residual analysis and centered leverage analysis [2, 7]. For instance, boxplots are commonly used to identify potential outliers in the distribution of a single-construct dataset, while scatterplots help identify outliers in the distribution of a multi-construct dataset. Within Exploratory Data Analysis (EDA), visualization serves not merely as a supportive tool but as an essential step. By extending human cognition to uncover patterns that might otherwise remain hidden, visualization offers the shortest path for analysts to gain insights into potential anomalies [28, 51]. Employing visualization prior to quantitative analysis helps overcome the limitations associated with relying exclusively on quantitative methods. Quantitative approaches for identifying outliers in multi-construct data, such as externally studentized analysis and centered leverage diagnostics, commonly assume specific statistical distributions of the residuals and linear relationships within the data [9, 48]. When these assumptions are violated—as is frequently the case in real-world data—quantitative outlier detection can become unreliable. Visualization can expose such violations by intuitively revealing nonlinearity (e.g., quadratic or cyclical trends) or unexpected clusters, thus laying a robust foundation for subsequent statistical validation.

Building on this foundational role, researchers have developed a wide range of visual strategies – going beyond basic tools like boxplots and scatterplots – tailored to more complex data types and analytical goals. For instance, outliers in temporal data can take many forms due to evolving patterns over time, making detection especially challenging without domain expertise. The MANDALA system by Suschnigg et al. [50] enables analysts to incorporate their domain knowledge to define what constitutes normal data points and interactively explore potential outliers using hexagonal binning and scatterplots. For high-dimensional data, Wilkinson et al. [54] introduced the HDOutliers algorithm, which applies random projection to reduce dimensionality and assigns a probability score to each data point, facilitating the visualization of outliers based on how unlikely a point is to belong to existing clusters. In the context of functional data, where each data object is an elastic curve, Xie et al. [58] proposed a geometrically motivated boxplot design that allows analysts to more easily detect componentwise outliers across different variabilities in a curve, such as shape, scale, and orientation. Thus, while some algorithmic or statistical approaches exist for outlier removal, visualization nevertheless can play a key role in outlier detection.

2.3 Confirmation Bias

Confirmation bias is a pervasive cognitive phenomenon in which individuals tend to seek out, interpret, and remember information in ways that reinforce their prior beliefs [42]. It can also be understood as a tendency to preserve, or a reluctance to abandon favored hypotheses [33, 34].

Confirmation bias can influence multiple stages of human hypothesis development [33]. For example, individuals may hold strong prior beliefs that persist even after receiving new evidence, consistent with a Bayesian updating process. They may also revise their confidence in a favored hypothesis insufficiently by downplaying the strength of contradictory evidence, or struggle to generate alternative hypotheses even when they sense limitations in their current one.

Confirmation bias can lead to a wide range of adverse consequences, varying in severity. On the less extreme end, it may cause researchers to favor evidence that supports their preferred theories, potentially compromising scientific objectivity [46]. More seriously, confirmation bias can contribute to political attitude polarization, reinforcing societal divisions [20]. It can also lead authorities to focus selectively on evidence that confirms a suspect’s guilt or innocence, thereby biasing legal investigations and prosecutions—sometimes resulting in wrongful convictions that permanently alter individuals’ lives [20, 39]. On the other hand, recent work also suggests that confirmation bias may also have adaptive functions. For instance, Peters (2022) introduced the reality-matching account, proposing that the persistence of confirmation bias is partly due to its social utility. According to this view, confirmation bias enables individuals to affect others’ beliefs and social structures to align with their own, thereby minimizing dissonance with the social world and enabling more effective social navigation [43].

The use of visualizations is not immune to the manifestation of confirmation bias. Prior visualization research has documented confirmation bias in cases where individuals interpret the same visualization in drastically different ways [41, 60], and have proposed mitigation strategies such as visually annotating contradictory data patterns to promote more balanced reasoning [15, 21]. In addition to visual content, textual elements in visualization designs – such as data summaries, captions, and title, also play a significant role in shaping the interpretation of the data and can also act as strong triggers of confirmation bias beyond visual contents in visualization designs [32, 35, 38].

Motivated by the influence of both visual and textual content on confirmation bias during the visualization interpretation, this work examines whether confirmation bias can also affect data cleaning processes supported by visualization – where users determine which portions of data are retained for further analysis and ultimately inform final decision-making.

3 STIMULUS DESIGN

We hypothesized that participants are likely to identify outliers when doing so produces a data pattern that aligns with their prior beliefs. To capture this behavior, we designed 2D scatterplot visualization stimuli that deliberately introduced ambiguity, such that removing subset of data points could reveal either belief-consistent or belief-inconsistent linear relationships. This design allows us to directly observe confirmation bias depending on participants’ data cleaning behavior, such as selectively identifying points whose exclusion results in a linear relationship between x and y that aligns with their beliefs.

3.1 Design Principles

To guide our stimulus design, we established four key principles.

Principle 1: The dataset should not contain a clearly dominant linear trend. A strongly dominant trend would encourage participants to converge on a common interpretation and reduce our ability to observe confirmation bias. To preserve ambiguity, we constrained the full dataset to have a correlation coefficient of 0 (absolute tolerance = 1×10^{-4}), ensuring that the scatterplot did not strongly suggest either a positive or a negative linear relationship.

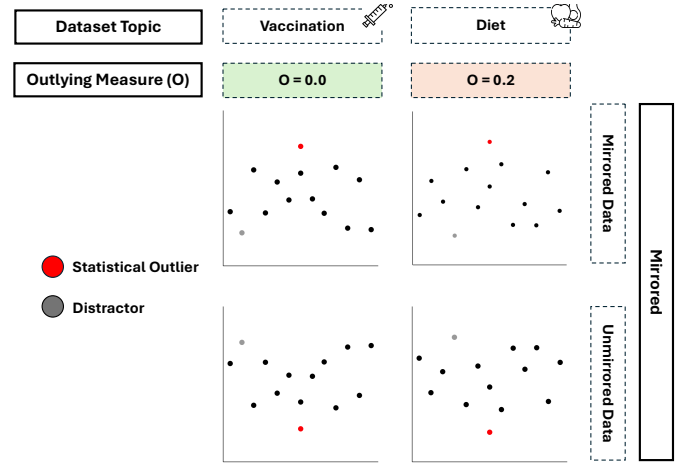


Figure 2: Main study stimulus conditions. *Dataset Topic*, *Outlying Measure*, and *Mirrored* were all between-subject factors.

Principle 2: The dataset should include legitimate outlier(s) to ensure participants can identify valid anomalies based solely on the underlying data pattern.

To differentiate belief-driven data removal from legitimate data-driven reasoning, the stimulus must afford participants a principled basis for identifying anomalies. Otherwise, it will be difficult to disentangle the mechanisms underlying participant behavior, as any removal of points could be interpreted as arbitrary or solely belief-driven. Thus, we had to make sure that each dataset contained at least one statistically legitimate outlier, such that participants could justify their data removal decisions based on the data distribution. We operationalized a **Statistical Outlier** as a data point with a studentized residual greater than 2, following established recommendations for two-dimensional data [9]. This model-based criterion allows us to compare removals that are statistically motivated against those that are primarily belief motivated.

Principle 3: Statistical Outlier(s) should not be too obvious. If the **Statistical Outlier** were too visually salient, participants would likely converge on the same data cleaning behavior of removing it and therefore undermining the intended ambiguity of the task. To balance statistical validity (Principle 2) with perceptual ambiguity (Principle 3), we constrained the **Statistical Outlier** to be only moderately extreme. We do this by specifying exactly one data point to have a studentized residual just above the conventional threshold (target = 2.1, absolute tolerance = 0.02). This ensures the presence of a statistically legitimate outlier that does not deviate conspicuously to dominate participant judgments.

To further reduce the perceptual distinctiveness of the **Statistical Outlier**, we introduced a visually similar but statistically non-outlying point (**Distractor**), defined as having a studentized residual in the range (1.5, 2). By placing the **Distractor** near the outlier threshold, we create a competing candidate for removal with comparable visual salience but weaker statistical justification. This design increases ambiguity in the outlier identification task, allowing us to better differentiate between statistically grounded and belief-driven cleaning decisions.

Principle 4: Removing only the statistical outlier should not change the dataset’s correlation. To fairly assess confirmation bias across participants, selecting the **Statistical Outlier** alone should not bias the data toward either direction. Otherwise, making a statistically-motivated decision to remove the **Statistical Outlier** could resemble a belief-motivated decision. We therefore constrained the correlation coefficient to remain 0 both for the full dataset and for the dataset after removing only the **Statistical Outlier** (absolute tolerance = 1×10^{-4}). Under this constraint, removing only the **Statistical Outlier** did not introduce directional bias and was therefore not treated as an instance of confirmation bias.

3.2 Stimulus Validation

To reduce unnecessary visual complexity while maintaining control over our design principles, we started with simpler scatterplots containing 14 data points. We then developed an R script to simulate candidate scatterplots that satisfied these four design principles (see Supplemental Materials).

In addition, to examine whether the visual salience of outliers influenced participants' outlier-identification behavior, we simulated datasets under two levels of the scagnostic *Outlying Measure*¹ defined by Wilkinson et al., 0 and 0.2 [56]. We selected these low and moderate values to vary the visual salience of outliers in the scatterplots while still maintaining the ambiguity required by Principle 1. We simulated 12 scatterplots using the script: six with an Outlying measure of 0 and six with an Outlying measure of 0.2. We then conducted a stimulus validation study to identify the scatterplot most likely to be perceived as having a slope of 0 for each *Outlying Measure*.

Participants. We recruited 40 participants from the crowdsourcing platform Prolific ($M_{age} = 40.65$, $SD_{age} = 9.75$; 21 Male, 17 Female, 1 Non-binary, 1 Agender). Each participant evaluated the slope of all 12 scatterplots in a random order.

Study procedure. After providing informed consent, participants completed a practice phase to become familiar with the interactive slider used to evaluate the slope of a scatterplot. The slider ranged from -1 to $+1$ in 0.1 increments, and the regression line on the scatterplot updated in real time to reflect the slope specified by the slider. In this phase, participants completed five easy trials with clearly distinguishable slopes. These practice scatterplots were simulated from a linear model, $y = \beta x + \epsilon$, with small Gaussian noise ($\epsilon \sim \mathcal{N}(0, \sigma^2)$, $\sigma = 0.2$) and fixed slopes $\beta \in \{-1, -0.5, 0, 0.5, 1\}$. During the practice phase, visual feedback indicated a correct response: the regression line turned green when the slider matched the target value. Participants were required to correctly set the slider for all five practice trials before proceeding to the formal evaluation, which consisted of 12 scatterplots.

Results. An ideal stimulus scatterplot is one that participants perceive as having a zero slope, so that in the formal study, we can measure confirmation bias as the extent to which the perceived slope shifts towards the direction of participants' prior beliefs. To reduce the risk of extreme slope interpretations, we selected stimuli based on the proportion of responses that fell within the $[-0.3, 0.3]$ range. Accordingly, for each *Outlying Measure* condition (0 vs. 0.2) [56], we chose the stimulus with the highest proportion of responses in this range, as shown in Figure 2.

4 MAIN STUDY

After selecting the stimuli, we conducted a pre-registered² experiment to assess the extent to which individuals exhibit confirmation bias in the identification and perception of outlier data points. To answer our research question, we adapted the 'Find Anomalies' task from the framework by Amar et al. [3], asking participants to visually identify outliers in a 2D scatterplot. Specifically, we asked participants to **Select** outlier data points (if they identified any) in a scatterplot by clicking to highlight them [61]. Our goal was to examine whether participants tend to **Select** data points that deviate from their expected linear trends, particularly in cases where the underlying relationship in the scatterplot is ambiguous. All studies have been reviewed and approved by Emory University's Institutional Review Board (IRB; STUDY00006185).

4.1 Belief Elicitation Methodology

Establishing users' prior beliefs is a critical step towards understanding how confirmation bias manifests. Thus, we need an elicitation method that allows users to map their internal beliefs to a fixed scale, which

¹The Outlying scagnostic ranges from 0 to 1 and captures how visually separated potential outlier(s) are from the rest of the scatterplot. A value of 0 indicates that no points are visually isolated from the main data cluster, while higher values indicate that one or more data points are increasingly separated from the rest of the data.

²<https://aspredicted.org/pm5b95.pdf>

also allows us to conduct quantitative analyses. To capture how users internally represent their beliefs in relation to our scatterplot stimulus, we adopt a visual belief elicitation technique inspired by Koonchanok et al. and Karduni et al., which allows participants to express both the direction and uncertainty of their prior belief as a linear relationship in two steps [31, 36].

- Step 1: Participants first interact with a line whose slope can be adjusted between -1 and $+1$ via a slider to indicate the perceived direction of the relationship between the two variables represented in the scatterplot.
- Step 2: While the line in step 1 remains visible, participants are then asked to indicate their uncertainty about this belief using a separate slider ranging from 0 to 1 as (σ), following the uncertainty elicitation approach of Koonchanok et al. [36].

The values from both sliders, including the direction (μ) and the uncertainty (σ), are subsequently used to parameterize the following linear model:

$$y = \beta x \quad (1)$$

$$\beta \sim \mathcal{N}(\mu, \sigma^2) \quad (2)$$

In Equation 1, β is a noisy slope of the linear relationship between x and y , modeled as a normally distributed variable. The mean of the distribution (μ) corresponds to the slope specified by participants in Step 1 as the direction of the prior belief, while the standard deviation is given by σ (Equation 2), which each participant provided in Step 2 to indicate the uncertainty of the prior belief. Following the approach of Koonchanok et al. [36], we use this model to animate a sequence of hypothetical outcome plots (HOPs) to help participants perceive the range of plausible data patterns implied by their indicated relationship directions and associated uncertainties. The HOPs were updated every 200 milliseconds. A depiction of three animation frames illustrating different levels of uncertainty in the expected relationship is shown in Figure 3.

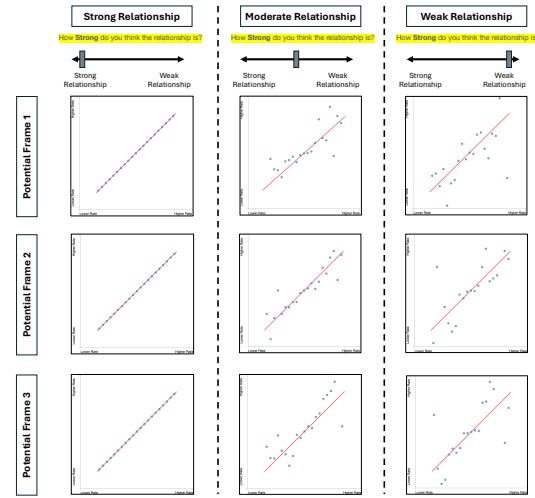


Figure 3: An illustration of potential frames from the HOPs animation during elicitation stage 2, under varying strengths of relationships.

4.2 Participants

We recruited 507 participants from Prolific ($M_{age} = 37.06$, $SD_{age} = 12.62$; 255 male, 235 female, 13 non-binary, 4 prefer not to disclose). Eligible participants were required to reside in the United States (per IRB), report English as their first language, and have a study approval rate above 95%. Individuals who had participated in the Stimuli Validation Study were not eligible for the main study. The sample size was determined through a power analysis conducted in G*Power [30], targeting 80% power at an alpha level of 0.05 across all hypotheses.

This analysis was based on pilot data from 67 participants collected under the same study setting as the main experiment.

4.3 Study Procedures

The experimental procedure is summarized in Figure 4. After completing the consent form, participants were guided through a tutorial that introduces the concept of outliers and familiarizes them with the visual interface used for outlier identification. The tutorial begins with a scatterplot containing 100 data points that depict a general positive correlation between high school students' height and weight, along with two outliers: one student with unusually short height and high weight, and another with unusually tall height and low weight. To avoid biasing the way participants identify outliers, these points are introduced as examples of outliers, which we defined as "exceptions to the relationship" between height and weight, following the definition of anomalies proposed by Amar et al. [3]. After this introduction, participants were asked to identify the two outliers by clicking to highlight them to complete the tutorial.

Following the tutorial, participants were asked to express their prior belief about the relationship between the two variables represented in the scatterplot, following the setup described in subsection 4.1. After the belief elicitation stage, participants proceeded to identify outliers in an interactive scatterplot. Participants were able to select and highlight potential outliers by clicking on individual data points and deselect them by clicking again. They were free to select as many or as few outliers as they deemed appropriate.

Next, participants evaluated the resulting visualization by specifying the perceived slope of the scatterplot (S_e) and answering a multiple-choice question designed to capture their rationale for selecting those outliers. To ensure data quality, two checks were included: a comprehension check immediately after the tutorial, which asks participants to identify the correct definition of outliers through a multiple-choice question; and an attention check following the evaluation stage, which states, "What is your favorite color? This is an attention check, please choose purple if you are paying attention." Participants who fail either the comprehension check or the attention check were excluded from the data analysis.

Lastly, participants proceeded to a final stage with two sensitivity checks designed to assess their slope perception. In these checks, participants evaluated the slope of two additional scatterplots without axis titles on a scale from -1 to $+1$. Both checks had a true regression slope of 0, so responses closer to 0 indicated more accurate slope perception. Check 1 assessed baseline slope perception using a scatterplot with an approximately horizontal trend. Because it served as a straightforward check of whether participants could identify a near-zero slope, participants whose responses fell outside the range of $[-0.3, 0.3]$ were excluded from the analysis. Check 2 was designed to assess participants' ability to perceive the slope of ambiguous scatterplots resembling those used in the main study. To avoid bias from reusing the exact main study stimulus, we presented a different scatterplot while preserving the participant's assigned *Outlying Measure* condition. Responses to Check 2 were used to account for individual differences in slope perception when analyzing confirmation bias.

4.4 Two Metrics for Measuring Confirmation Bias

One manifestation of confirmation bias is the tendency to interpret ambiguous evidence as supporting evidence [13, 17, 33]. Building on this idea, we aim to measure confirmation bias by examining the extent to which individuals either:

1. Interpret an ambiguous scatterplot as depicting a correlation with a direction aligned with their prior beliefs (hereafter, *perception bias*), or
2. Alter an ambiguous scatterplot to reflect such a correlation through outlier identification (hereafter, *removal bias*).

As depicted in Figure 1, we distinguish between two types of confirmation bias: *perception bias* reflects participants' tendency to interpret an ambiguous visualization as exhibiting belief-consistent correlation; *removal bias* refers to their outlier identification behaviors, specifically,

whether participants classify data points that contradict their prior beliefs as outliers (*removal bias* = 0 when participants do not remove any data points).

To facilitate quantification of both *perception bias* and *removal bias* as the magnitude of the slope in the direction implied by the priors, we incorporated both Principle 1 and Principle 4 as described in subsection 3.1.

- *perception bias* (CB_p): We measure *perception bias* as the difference between a participant's estimated regression line slope (s_e) in the resultant scatterplot and the true underlying slope of these resultant visualizations (0), taken in the direction of the participant's prior beliefs. Accordingly, CB_p is defined as:

$$CB_p = \pm |s_e| \quad (3)$$

In this equation, s_e represents the participant's estimated slope of the regression line in the scatterplot stimulus after their outlier selections are removed. The value of CB_p is positive (+) when the sign of s_e matches the sign of the participant's prior belief slope elicited in Stage 1, and negative (−) when the two signs diverge. For example, if a participant holds a prior belief in a positive relationship and evaluates the slope of the post-removal scatterplot as 0.2, then CB_p would equal 0.2.

- *removal bias* (CB_r): We measure *removal bias* as the difference between the slope of the regression line in the scatterplot after participants' selected outliers are removed (s_o), and the slope of the regression line in the scatterplot with only the statistical outlier removed (0) taken in the direction of participants' prior belief. Since no data points other than the **Statistical Outlier** should be identified as outliers based on the underlying data structure of our stimuli, we measure *removal bias* as the extent to which participants modified the original neutral stimulus (slope = 0) into a scatterplot depicting a correlation aligned with their prior beliefs:

$$CB_r = \pm |s_o| \quad (4)$$

In this equation, s_o indicates the underlying slope of the scatterplot after the removal of participants outlier selections. The value of CB_r is positive (+) when the sign of s_o matches the sign of the participant's prior belief slope elicited in Stage 1, and negative (−) when the two signs diverge. For instance, if a participant holds a prior belief in a positive relationship and removes two points whose removal shifts the slope of the scatterplot from 0 to 0.2, then CB_r would be 0.2. For participants who choose not to remove any data point, we can still measure confirmation bias from interpretations (*perception bias*) using Equation 3.

4.5 Hypotheses

Guided by prior visualization and psychology research, we formulated hypotheses for *perception bias* and *removal bias* – two measures of confirmation bias in our study – to test whether participants interpret the final visualization or remove data points in ways that align with their prior beliefs. Prior research has shown that when the trend in a visualization is ambiguous without explanatory text, readers often resolve that ambiguity in a belief-consistent direction, with stronger priors amplifying the effect [38]. Building on this evidence, we propose the following hypotheses for *perception bias*:

- **H1:** Participants will exhibit *perception bias*, perceiving the slope of the scatterplots as consistent with their prior beliefs.
- **H2:** The level of *perception bias* will increase as prior belief strength increases.

Likewise, prior work by Xiong et al. exemplified that viewers may selectively include or discount outliers depending on their opinions

²For **H3**, under our study design, the *removal bias* score was defined as zero when no data points were removed. However, because the regression model did not enforce this rule, the estimated marginal means for no removal remained centered around zero rather than fixed at zero.

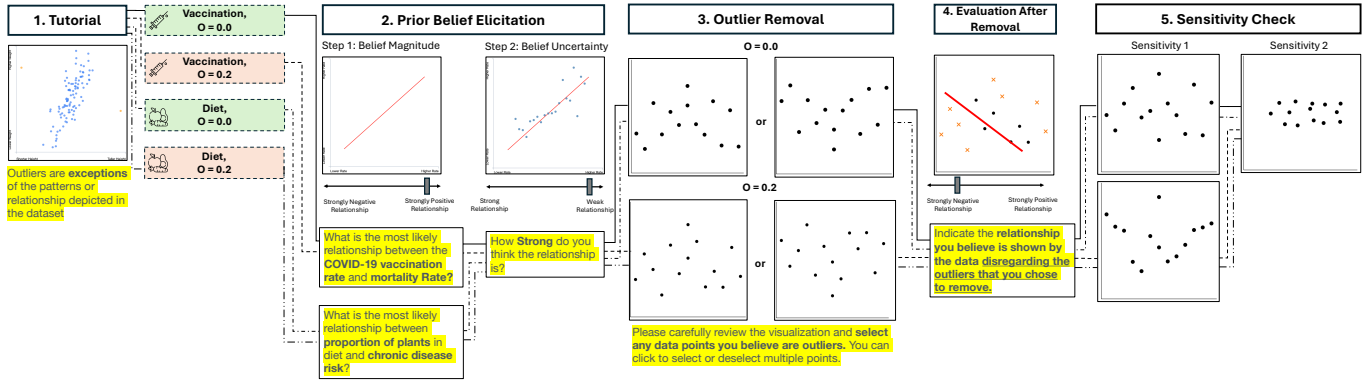


Figure 4: Main study procedure: each participant completed one trial, in which they removed outliers from a scatterplot assigned to one of four visualization conditions defined by the between-subject conditions *Topic* and *Outlying Measure*. In addition, the scatterplot was randomly displayed in either its original or Mirrored form.

about the scatterplot, in ways that support favored outcomes [6, 59]. Similarly, psychology research on scatterplot reasoning has further identified that outlier judgments and perceived regression-line slope are more accurate when outliers lie farther from the regression-line [16]. Motivated by these findings, we propose the following hypotheses for *removal bias*:

- **H3:** Participants will exhibit *removal bias*, removing data points in ways that make the resulting visualization consistent with their prior beliefs.
- **H4:** The level of *removal bias* will increase as prior belief strength increases.
- **H5:** The level of *removal bias* will increase as outliers become less visually salient.

4.6 Materials

- **Topic:** To test H1–5 across a broader range of pre-existing belief strengths, we selected one polarizing topic [COVID-19 Vaccination] and one more neutral topic [Diet Choices], both validated in prior visualization research [38]:
 - **[COVID-19 Vaccination]:** Participants reason with a scatterplot depicting the relationship between county-level COVID-19 vaccination rate and mortality rate in the United States.
 - **[Diet Choice]:** Participants reason with a scatterplot showing the relationship between the proportion of plant food in one’s diet and chronic disease risk.
- **Outlying Measure:** As mentioned in section 3, to test H5, we manipulated the scagnostic Outlying measure to vary the visual salience of the statistical outlier:
 - **[Small Outlyingness (0)]:** The outlier is not visually obvious
 - **[Middle Outlyingness (0.2)]:** The outlier is moderately salient-noticeable but not obvious
- **Visual Stimuli:** We selected two scatter plots – one with a [Small Outlying] measure and one with a [Middle Outlying] measure – based on results from the stimuli validation study (see subsection 3.2).

4.7 Conditions

As depicted in Figure 2, both *Topic* and *Outlying Measure* are between-subject conditions; each participant was randomly assigned to one of four scatterplots: the *Topic* was either [COVID-19 Vaccination] or [Diet Choice], and the *Outlying Measure* was either [Small Outlyingness] or [Middle Outlyingness]. Additionally, we included mirrored versions of each scatterplot (vertically flipped variants) to counterbalance any orientation-driven tendency to infer a direction trend from noise, reducing systematic bias unrelated to our experimental manipulations.

5 RESULTS

5.1 Measures

We first used a one-sample t-test to examine whether *perception bias* (H1) was exhibited overall. We then fit two fixed-effects linear regression models: one for *perception bias* to test H2, and one for *removal bias* to test H3, H4, and H5. In each model, the dependent variable was computed using Equation 3 or Equation 4, respectively. Both models included independent variables directly related to our hypotheses, along with additional study-design factors and participant-level covariates:

- **Belief Magnitude:** the absolute value of the slope selected by participants in prior elicitation Step 1 (ranging from 0 to 1 in increments of 0.1)
- **Belief Certainty:** the inverse of the uncertainty level reported in prior elicitation Step 2, computed as $1 - \text{Reported Uncertainty}$ (ranging from 0 to 1 in increments of 0.05)
- **Dataset Topic:** topic of the dataset variables (COVID-19 vaccination or diet choice).
- **Outlying Measure:** the scagnostic outlying measure of the dataset [56] (0 or 0.2)
- **Mirrored:** whether the dataset was reflected across the y-axis (“yes” or “no”)
- **Sensitivity Score:** the absolute value of the slope in Check 2 (ranging from 0 to 1 in increments of 0.1)
- **Data Point Removed:** whether the participant removed any data point (“yes” or “no”)

The modeling results for both *perception bias* and *removal bias* are summarized in Table 1, and Figure 5 summarizes our findings for the five hypotheses.

5.2 Analysis for H1

To examine whether participants exhibited *perception bias* overall, we conducted a two-sided one-sample t-test comparing *perception bias* scores against 0. A mean significantly greater than 0 indicates that, on average, participants perceived the data as more consistent with their prior beliefs.

The mean *perception bias* score was significantly greater than 0 ($M = 0.178$, 95% CI [0.147, 0.209]), $t(506) = 11.38$, $p < .001$. This indicates that, overall, participants interpreted a neutral scatterplot as showing a correlation aligned with their prior beliefs, thereby **supporting H1**.

5.3 Analysis for H2

We fitted a linear model with the seven independent variables described in subsection 5.1 to examine how prior belief strength influenced the exhibited level of *perception bias*. In particular, we examined the main effect of Belief Magnitude on *perception bias* scores. As shown in Figure 5 and Table 1, results for H2 indicate that Belief Magnitude

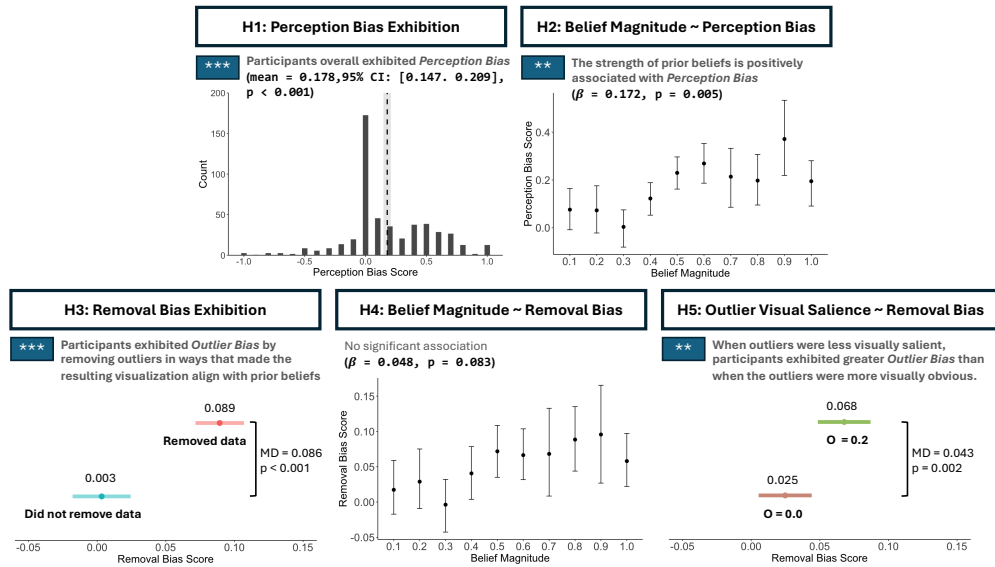


Figure 5: Main study results. For **H1**, the bar chart shows the distribution of participants' *perception bias* scores. The dashed vertical line marks the mean *perception bias* score, and the gray shaded region indicates the confidence interval. The results for **H2** and **H4** are shown as error bars for each Belief Magnitude level. The results for **H3** and **H5** are presented as confidence interval plots, which show the estimated marginal means and mean differences (MD) of the *removal bias* score with confidence intervals for Data Point Removed and Outlying Measure.³

was significantly associated with the level of *perception bias* ($p = 0.005$). Specifically, a one-unit increase in Belief Magnitude was associated with a 0.172 shift in slope interpretation toward participants' prior beliefs, thereby **supporting H2**.

5.4 Analysis for H3, H4, and H5

We fitted another linear model on *removal bias* using the same seven independent variables. To test **H3**, we examined the main effect of Data Point Removed to assess whether choosing to remove at least one data point was associated with higher *removal bias* scores, indicating that the resulting visualization was more aligned with participants' prior beliefs. To test **H4**, we examined the main effect of Belief Magnitude on *removal bias* scores. To test **H5**, we examined the main effect of Outlying Measure on *removal bias* scores to determine whether the visual saliency of outliers affected the exhibited level of *removal bias*. Specifically, Outlying Measure had two levels, 0 and 0.2: a value of 0 indicates that the **Statistical Outlier** and **Distractor** were less visually salient in the scatterplot, whereas a value of 0.2 indicates that they were slightly more salient.

According to the model fitting results in Table 1, **H3** was supported by a significant main effect of Data Point Removed ($p < 0.001$). Specifically, choosing to remove at least one data point was associated with an approximately 0.1 shift in the slope of the finalized visualization toward the direction of participants' prior beliefs.

For **H4**, Belief Magnitude showed a marginal but not statistically significant association with *removal bias* ($p = 0.083$); therefore, **H4 was not supported**.

For **H5**, we identified a significant main effect of Outlying Measure on *removal bias* scores ($p = 0.002$). Specifically, when participants viewed scatterplots in which the **Statistical Outlier** and **Distractor** were less visually salient, their outlier removals shifted the slope of the finalized visualization by about 0.1 toward their prior beliefs. These results **support H5**.

6 DISCUSSION

6.1 The Impact of Outlier Removal on Perception Bias

Our results suggested that both Belief Magnitude and Data Point Removed were positively associated with *perception bias* (see Table 1). Conceptually, however, we see outlier removal as an intermediary step that allows participants to transform the evidence they later evaluate.

As participants removed the data points, they changed the dataset in a way that potentially better matched what they already believed. As a result, we might see their prior beliefs show up more pronounced in how they interpret the final visualization after the outlier removal. To examine this possibility, we conducted an exploratory analysis testing whether the association between Belief Magnitude and *perception bias* varied as a function of whether participants removed any data points.

We added an interaction term between Belief Magnitude and Datapoint Removed to the *perception bias* model described in [subsection 5.1](#) and refit the model. The interaction was statistically significant ($p = 0.01$). Specifically, among participants who removed at least one data point, Belief Magnitude was positively associated with *perception bias* ($\beta = 0.291$, $p = 0.01$). In contrast, among participants who did not remove any data points, Belief Magnitude was not significantly associated with *perception bias* ($\beta = -0.004$, $p = 0.969$).

As illustrated in Figure 6, our exploratory analysis suggested that stronger prior beliefs were positively associated with more belief-consistent interpretations only among participants who removed at least one data point. Participants who did not remove any data points interpreted the resulting visualizations in ways that were less aligned with their prior beliefs, even when those beliefs were strong.

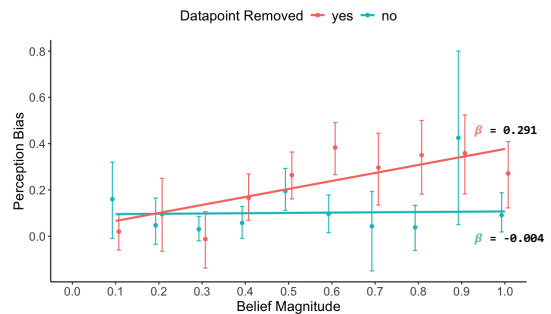


Figure 6: The association between Belief Magnitude and *perception bias* among participants who removed and did not remove any data points.

These findings suggests several implications for the data cleaning

Table 1: fixed-effects linear regression results for *perception bias* and *removal bias*. Because each predictor in the main models has one degree of freedom, the coefficient-test p -values and Type II ANOVA p -values are identical; therefore, we report a single p -value for each predictor.

Independent Variable	<i>perception bias</i>				<i>removal bias</i>			
	Coefficient	Std. Error	t-value	p	Coefficient	Std. Error	t-value	p
Belief Magnitude	0.172	0.060	2.853	.005 **	0.048	0.028	1.735	.083
Belief Certainty	0.007	0.073	0.092	.927	0.052	0.033	1.542	.124
Dataset Topic (vaccination)	0.003	0.030	0.087	.930	0.016	0.014	1.175	.241
Outlying Measure (0.2)	0.055	0.030	1.840	.066	0.043	0.014	3.113	.002 **
Mirrored (yes)	0.021	0.030	0.688	.492	-0.020	0.014	-1.441	.150
Sensitivity Score	0.307	0.062	4.952	< .001 ***	0.052	0.028	1.825	.069
Data Point Removed (yes)	0.101	0.031	3.268	.001 **	0.086	0.014	6.078	< .001 ***

practice and the design of future visualization systems that support visual data cleaning. In particular, our findings reinforce the importance of specifying outlier-identification and handling criteria before inspecting whether particular removal decisions could produce belief-consistent patterns. When such *a priori* specification is not possible, such as in early-stage data-cleaning workflows where analysts must first inspect the data to understand what kinds of data-quality issues are present, future systems could support side-by-side views of data before and after removal, showing how correlation between variables changes after removal, or encourage metacognitive reflection by explicitly asking users about their confidence and justification for each removal decision [14]. Beyond reducing bias introduced by removal in downstream analysis, such interventions may also decrease the likelihood that users perceive data inaccurately to align with their prior beliefs. We approach these implications not as prescriptive guidelines but as a way to envision how future systems could promote reflective outlier handling processes, since our study reveals how confirmation bias shapes outlier-removal but does not directly evaluate specific interventions to mitigate this bias.

6.2 Slope Interpretation Accuracy in Confirmation Bias

Prior work has shown that visual estimations of slopes in scatterplots can differ from statistical trends in predictable ways: viewers tend to discount the influence of outliers during slope estimation [11, 19], and may overestimate slope steepness when estimating or validating linear slopes by eye [10]. In our context, this might raise an alternative explanation: **the observed *perception bias* may reflect general variability in human regression-by-eye judgments rather than confirmation bias**. In our stimulus design, we constrained both the full dataset and the dataset after removing only the **Statistical Outlier** to have a correlation of zero. Thus, we reduced the possibility that deviations in *perception bias* above or below zero can be attributed to the presence of the **Statistical Outlier**. Additionally, if *perception bias* only reflected random variability in trend estimation, its mean value should be centered around zero. However, as reported in subsection 5.2, the mean value of *perception bias* was significantly greater than zero ($M = 0.178$, $p < .001$), suggesting that participants' estimations were generally biased toward prior beliefs. This pattern is more consistent with confirmation bias than with estimation variability alone.

To evaluate whether variability in slope estimation could account for the observed bias patterns, we next examine whether individual differences in slope estimation accuracy are associated with *perception bias* and *removal bias*. We hypothesized that participants with low Sensitivity Score would exhibit less bias, such that participants who more accurately interpreted a context-free, zero-correlation scatterplot would show lower *perception bias* and *removal bias*.

Consistent with this hypotheses, we found that Sensitivity Score was significantly associated with *perception bias* ($p < 0.001$) and showed a marginal positive association with *removal bias* ($p = 0.069$). These results suggest that (i) participants with worse Sensitivity Score were more likely to interpret the data in ways that are consistent with their prior beliefs (*perception bias*), and (ii) they were more likely to remove data points in ways that increased

consistency with those beliefs (*removal bias*).

One possible explanation is that participants differed in their threshold for perceiving structure in noisy or ambiguous data, leading some to infer linear trends even when none are present. From a signal detection perspective, participants with higher Sensitivity Score may have lower perceptual precision when estimating correlation, resulting in greater perceptual uncertainty. This increased uncertainty may, in turn, promote a more liberal decision criterion during outlier identification. Under such conditions of uncertainty, reader judgments may rely more heavily on prior beliefs, increasing the likelihood of belief-consistent interpretations. This pattern aligns with prior work showing that people make more inferences from weak signals, such that when evidence is ambiguous, they rely more on prior beliefs and may overinterpret noise as signal [4].

Alternatively, participants' interpretations of the study stimuli after outlier removal might have been subject to an anchoring effect on their responses in the second sensitivity check (Check 2), leading them to report slope evaluations similar to their earlier judgments [5, 24, 52]. To explore this possibility, we conducted an additional Pearson correlation test between participants' post-removal slope evaluation (S_e) and their slope evaluations of the scatterplot in the second sensitivity check [8]. The correlation was weakly positive, though statistically significant ($r = 0.148$, $p = 0.001$). This suggests that participants' post-removal interpretations might have slightly anchored their evaluations in the second sensitivity check.

Future work should compare participants' slope sensitivity both before and after primary study stimuli. Understanding this possibility is important because people often encounter visualizations sequentially in practice. From understanding political trends through data journalism to monitoring product performance from visualization dashboards, sequential decision-making using data visualization might also introduce anchoring bias, and thus, influencing how individuals perceive the world and make decisions based on these visualizations [23].

6.3 Reasons Behind Data Point Removal

As described in subsection 4.2, after completing the outlier-removal task, participants were asked to report their reason for removing or not removing outliers through a multiple-choice question with five options:

- R1: Chosen as **exceptions to visual patterns**.
- R2: Chosen because they do **not fit the expected relationship**.
- R3: Chosen for a reason **not listed above** (specify).
- R4: No points chosen; saw no **exceptions to visual patterns**.
- R5: No points chosen for a reason **not listed above** (specify).

Among participants who provided a valid reason for their removal decision (e.g., they reported a reason consistent with their action, whether it be removing or not removing data), about 59% ($n = 293$) removed at least one data point. Among these participants, the majority reported removing points they perceived as visual exceptions (R1, $n = 213$), whereas a smaller portion acknowledged removing data points that were inconsistent with their prior beliefs (R2, $n = 65$). However, participants who reported that the removed data points were visual exceptions still, overall, removed data points in ways that made the remaining data more aligned with their prior beliefs.

To examine this, we conducted a one-sample t-test on the **removal bias** scores of participants who selected R1, and we found that the mean **removal bias** was significantly greater than zero ($M = 0.086$, 95% CI [0.057, 0.115], $t(212) = 5.908$, $p < 0.001$). This suggests that even when participants described their decisions as selecting visual exceptions, their outlier removals were still associated with belief-consistent changes in the resulting data.

Among participants who did not remove any data points ($n = 205$), most reported doing so because they did not perceive any visual exceptions (R4, $n = 183$). However, their interpretations of the resulting visualization still tended to align with their prior beliefs. This suggests that confirmation bias may operate implicitly, even when users frame their actions as data-driven, and self-reported reasoning alone may be insufficient to capture bias during data cleaning.

Justification for Removing Data Points: We identified three themes among 15 specified reasons for removing data points (R3). The most common was **Visual Exception / Far from Other Points** ($n = 10$). While these responses broadly aligned with R1, participants articulated different forms of visual exceptions in their explanations. P71 described the selected point as “*really out of place*,” and P458 noted that the selected points were the *furthest away from the other points*.”

The second most frequent theme was **Uncertainty** ($n = 3$), where participants expressed hesitation in their removal decisions due to the ambiguity of the data. For instance, P258 explained, “*The pattern is not strong. It will be an outlier for either positive or negative relationship, except if the trend line is parallel to the x-axis*,” while P344 expressed this uncertainty more directly in the response, “*I picked the extremes though I didn’t find them to be so extreme and I wasn’t confident*.”

Finally, we identified responses that restated **Inconsistency with Expected Relationships** ($n = 2$), a rationale already captured by R2. For example, P191 wrote, “*they do not fit the reasonable relationship between proportion of plant foods in diet and chronic disease risk*,” and, notably, P92 described the decision as a combination of R1 and R2: “*A combination of exception to the visual pattern & do not fit the reasonable relationship*.” Overall, self-reported reasons for removing data points still centered on R1 and R2, with only a small number of participants expressing uncertainty about their removals.

Justification for Retaining Data Points: We identified five themes among the 22 specified reasons for retaining all data points (R5). Eleven participants reiterated R4, commenting **No Clear Pattern** in the dataset. For example, P34 observed that “*these points show that there is no clear relationship*.” Participants also retained data due to **Insufficient Sample Size** ($n = 4$). For instance, P222 commented, “*I believe there wasn’t enough data to clearly identify outliers*,” whereas P194 linked sample size to the credibility of the data: “*The sample of participants did not seem large ... That would ...alter the credibility of the data*.”

Additionally, a **Lack of Confidence in Candidate Outliers** ($n = 3$) and **Uncertainty in Judgment Ability** ($n = 2$) also led to data retention. P6 explained “*There were a couple outliers from the overall trend but they weren’t far enough away ...to disregard them*.”

Interestingly, only two participants were **Aware of Potential Bias** introduced by removal and thus decided to keep all data points. P33 insisted “*I believe that everyone’s data should matter, not just ones that fit the visual pattern*.” P56 stated more explicitly about bias from expectations, “*I did not see a visual pattern and I did not want to color the results with any biases I may have had about expected results*.”

6.4 Limitations and Future Work

Our study focuses on outlier identification through visual inspection. In practice, analysts often complement visual inspection with statistical tools, automated detection methods, and domain knowledge when deciding whether to remove data points [2, 37]. However, these workflows are susceptible to confirmation bias. For example, analysts can choose which detection method to use, adjust outlier-detection thresholds, decide whether to retain, transform, or remove potential outliers, or selectively invoke domain knowledge to justify particular handling decisions. They may also reference preliminary visualizations to iteratively manipulate statistical rules or data transformations until the

output aligns with their expectations. Our findings are therefore most relevant to these human-in-the-loop moments, where prior beliefs may shape how analysts identify outliers and the consequences of different outlier-handling decisions. Future work should examine whether similar belief-consistent patterns emerge when statistical or automated methods are used to detect and handle outliers.

Relatedly, the criterion of our participant selection did not include advanced statistical training or domain expertise on the dataset topics. Analysts with stronger statistical expertise may rely more heavily on formal diagnostics, such as studentized residuals, or other statistical criteria. Similarly, analysts with domain expertise could be better positioned to judge whether a data point is plausible, erroneous, or meaningful based on contextual knowledge. Future work should therefore examine whether statistical training and domain expertise moderate the influence of prior beliefs on outlier-removal decisions.

Another limitation is that, although the observed effects of **perception bias** and **removal bias** were statistically significant, their raw magnitudes were modest. Specifically, the overall **perception bias** effect exhibits medium effect size (Cohen’s $d \approx 0.51$), and the effect of removing data point(s) on **removal bias** also corresponds to medium effect size ($\eta^2 \approx 0.07$). Rather than suggesting that prior beliefs overwhelmingly dictate visual outlier removal, our findings suggest that prior beliefs still introduce a small but systematic, belief-consistent shift in how people interpret and modify an ambiguous scatterplot. These minor shifts can distort downstream analyses, model estimates, and data-driven decisions.

Additionally, our datasets only included 14 data points. Future work can also explore a range of datasets beyond what we tested. Existing work suggests that increasing the number of points in a dataset can reduce perceptual accuracy in correlation estimation tasks and outlier detection, especially when data patterns are not visually salient [53]. Future work can explore whether these effects persist at larger scales, or when participants are told that the data represent a sample rather than the full dataset. Beyond dataset size, our datasets were deliberately simplified to isolate confirmation bias in outlier-removal decisions: our design principles constrained each dataset to have zero correlation and to include a **Statistical Outlier** whose removal would leave the correlation unchanged. Therefore, our findings may not generalize to real-world data visualizations that contain stronger visible trends or use text annotations to reinforce a particular argument.

Furthermore, we only tested two-dimensional scatterplots. Real-world datasets are often much messier and many contain hundreds or even thousands of dimensions [44]. As a result, the bias patterns we observed may not fully generalize to more complex analytic settings with higher-dimensional datasets where outlier detection depends on multivariate visualizations such as parallel coordinates and computational methodologies [1, 12, 54, 55]. Future work could extend our paradigm to multivariate analysis settings and examine how these more complex beliefs influence analysis [57].

Finally, because prior belief elicitation was required to measure confirmation bias, it may have increased the salience of participants’ priors and strengthened belief-consistent slope interpretations and removal decisions. Future work could test whether **perception bias** and **removal bias** change when priors are elicited in a separate session or separated from the main task by a filler task.

7 CONCLUSION

Through a pre-registered study, we showed that participants exhibit two manifestations of confirmation bias in visual outlier identification for scatterplots: outlier removal bias and post-removal perception bias, particularly when potential outliers were less visually salient. When cleaning data points, participants systematically brought the post-removal slope closer to their prior beliefs; furthermore, those with stronger initial beliefs were significantly more likely to interpret the remaining data in a belief-consistent manner. These results suggest that prior expectations heavily dictate how people separate noise from signal to generate visual takeaways from data.

ACKNOWLEDGMENTS

This work was supported by the National Science Foundation under grants IIS-2311574. We also thank Mengyu Chen and Thomas Davidson for their extensive feedback.

REFERENCES

- [1] C. C. Aggarwal and P. S. Yu. Outlier detection for high dimensional data. In *Proceedings of the 2001 ACM SIGMOD international conference on Management of data*, pp. 37–46, 2001. doi: [10.1145/375663.375668](https://doi.org/10.1145/375663.375668) 9
- [2] H. Aguinis, R. K. Gottfredson, and H. Joo. Best-practice recommendations for defining, identifying, and handling outliers. *Organizational research methods*, 16(2):270–301, 2013. doi: [10.1177/1094428112470848](https://doi.org/10.1177/1094428112470848) 1, 2, 9
- [3] R. Amar, J. Eagan, and J. Stasko. Low-level components of analytic activity in information visualization. In *IEEE Symposium on Information Visualization*, 2005. *INFOVIS 2005.*, pp. 111–117. IEEE, 2005. doi: [10.1109/INFVIS.2005.1532136](https://doi.org/10.1109/INFVIS.2005.1532136) 2, 4, 5
- [4] N. Augenblick, E. Lazarus, and M. Thaler. Overinference from weak signals and underinference from strong signals. *The Quarterly Journal of Economics*, 140(1):335–401, 2025. doi: [10.1093/qje/qjae032](https://doi.org/10.1093/qje/qjae032) 8
- [5] Š. Bahník, B. Englich, and F. Strack. Anchoring effect. In *Cognitive illusions*, pp. 223–241. Psychology Press, 2016. doi: [10.4324/9781315696935](https://doi.org/10.4324/9781315696935) 8
- [6] E. Balcetiš and D. Dunning. See what you want to see: motivational influences on visual perception. *Journal of personality and social psychology*, 91(4):612, 2006. doi: [10.1037/0022-3514.91.4.612](https://doi.org/10.1037/0022-3514.91.4.612) 6
- [7] D. A. Belsley, E. Kuh, and R. E. Welsch. *Regression diagnostics: Identifying influential data and sources of collinearity*. John Wiley & Sons, 2005. doi: [10.1002/0471725153](https://doi.org/10.1002/0471725153) 2
- [8] J. Benesty, J. Chen, Y. Huang, and I. Cohen. Pearson correlation coefficient. In *Noise reduction in speech processing*, pp. 1–4. Springer, 2009. doi: [10.1007/978-3-642-00296-0_5](https://doi.org/10.1007/978-3-642-00296-0_5) 8
- [9] D. Blatná. Outliers in regression. *Trutnov*, 30:1–6, 2006. 2, 3
- [10] D. Braun, R. Chang, M. Gleicher, and T. von Landesberger. Beware of validation by eye: Visual validation of linear trends in scatterplots. *IEEE transactions on visualization and computer graphics*, 31(1):787–797, 2024. doi: [10.1109/TVCG.2024.3456305](https://doi.org/10.1109/TVCG.2024.3456305) 8
- [11] D. Braun, D. Eberle, R. Chang, M. Gleicher, and T. von Landesberger. Judging Lines, Ignoring Outliers? Influence of Outliers for Visual Validation versus Visual Estimation of Linear Trends. In D. Archambault, I. Nabney, and J. Peltonen, eds., *Machine Learning Methods in Visualisation for Big Data*. The Eurographics Association, 2026. doi: [10.2312/mlvis.20261004](https://doi.org/10.2312/mlvis.20261004) 8
- [12] R. J. Campello, D. Moulavi, A. Zimek, and J. Sander. Hierarchical density estimates for data clustering, visualization, and outlier detection. *ACM Transactions on Knowledge Discovery from Data (TKDD)*, 10(1):1–51, 2015. doi: [10.1145/2733381](https://doi.org/10.1145/2733381) 9
- [13] S. D. Charman, M. Kavetski, and D. H. Mueller. Cognitive bias in the legal system: Police officers evaluate ambiguous evidence in a belief-consistent manner. *Journal of Applied Research in Memory and Cognition*, 6(2):193–202, 2017. doi: [10.1016/j.jarmac.2017.02.001](https://doi.org/10.1016/j.jarmac.2017.02.001) 5
- [14] M. Chen, A. Yang, S. Min, K. A. Hamilton, and E. Wall. A novel lens on metacognition in visualization. In *Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems*, pp. 1–16, 2025. doi: [10.1145/3706598.3714400](https://doi.org/10.1145/3706598.3714400) 8
- [15] T. Chuanromanee and R. Metoyer. A crowdsourced study of visual strategies for mitigating confirmation bias. In *2022 IEEE Symposium on Visual Languages and Human-Centric Computing (VL/HCC)*, pp. 1–6. IEEE, 2022. doi: [10.1109/VL/HCC53370.2022.9833151](https://doi.org/10.1109/VL/HCC53370.2022.9833151) 3
- [16] L. Ciccione, G. Dehaene, and S. Dehaene. Outlier detection and rejection in scatterplots: Do outliers influence intuitive statistical judgments? *Journal of Experimental Psychology: Human Perception and Performance*, 49(1):129, 2023. doi: [10.1037/xhp0001065](https://doi.org/10.1037/xhp0001065) 6
- [17] M. B. Cook and H. S. Smallman. Human factors of the confirmation bias in intelligence analysis: Decision support from graphical evidence landscapes. *Human Factors*, 50(5):745–754, 2008. doi: [10.1518/001872008X354183](https://doi.org/10.1518/001872008X354183) 5
- [18] R. D. Cook. Detection of influential observation in linear regression. *Technometrics*, 19(1):15–18, 1977. doi: [10.1080/00401706.1977.10489493](https://doi.org/10.1080/00401706.1977.10489493) 2
- [19] M. Correll and J. Heer. Regression by eye: Estimating trends in bivariate visualizations. In *Proceedings of the 2017 CHI Conference on Human Factors in Computing Systems*, pp. 1387–1396, 2017. doi: [10.1145/3025453.3025922](https://doi.org/10.1145/3025453.3025922) 8
- [20] M. Del Vicario, A. Scala, G. Caldarelli, H. E. Stanley, and W. Quattrociocchi. Modeling confirmation bias and polarization. *Scientific reports*, 7(1):40391, 2017. doi: [10.1038/srep40391](https://doi.org/10.1038/srep40391) 3
- [21] R. Delpish, S. Jiang, L. Davis, and K. Odubela. A visual analytics approach to combat confirmation bias for a local food bank. In *Advances in Human Error, Reliability, Resilience, and Performance: Proceedings of the AHFE 2018 International Conference on Human Error, Reliability, Resilience, and Performance*, July 21–25, 2018, Loews Sapphire Falls Resort at Universal Studios, Orlando, Florida, USA 9, pp. 13–23. Springer, 2019. doi: [10.1007/978-3-319-94391-6_2](https://doi.org/10.1007/978-3-319-94391-6_2) 3
- [22] E. Dimara and J. Stasko. A critical reflection on visualization research: Where do decision making tasks hide? *IEEE Transactions on Visualization and Computer Graphics*, 28(1):1128–1138, 2021. doi: [10.1109/TVCG.2021.3114813](https://doi.org/10.1109/TVCG.2021.3114813) 1
- [23] J. M. Echterhoff, M. Yarmand, and J. McAuley. Ai-moderated decision-making: Capturing and balancing anchoring bias in sequential decision tasks. In *Proceedings of the 2022 CHI Conference on Human Factors in Computing Systems*, pp. 1–9, 2022. doi: [10.1145/3491102.3517443](https://doi.org/10.1145/3491102.3517443) 8
- [24] A. Furnham and H. C. Boo. A literature review of the anchoring effect. *The journal of socio-economics*, 40(1):35–42, 2011. doi: [10.1016/j.soccec.2010.10.008](https://doi.org/10.1016/j.soccec.2010.10.008) 8
- [25] D. Ghosh and A. Vogt. Outliers: An evaluation of methodologies. In *Joint statistical meetings*, vol. 12, pp. 3455–3460, 2012. 1
- [26] M. Guo, Y. Wang, Q. Yang, R. Li, Y. Zhao, C. Li et al. Normal workflow and key strategies for data cleaning toward real-world data. *Interactive Journal of Medical Research*, 12(1):e44310, 2023. doi: [10.2196/44310](https://doi.org/10.2196/44310) 1
- [27] D. M. Hawkins. *Identification of outliers*, vol. 11. Springer, 1980. doi: [10.1007/978-94-015-3994-4](https://doi.org/10.1007/978-94-015-3994-4) 1
- [28] N. A. Heckert, J. J. Filliben, C. M. Croarkin, B. Hembree, W. F. Guthrie, P. Tobias et al. Handbook 151: Nist/seamtech e-handbook of statistical methods. 2002. doi: [10.18434/M32189](https://doi.org/10.18434/M32189) 2
- [29] S. Kandel, A. Paepcke, J. Hellerstein, and J. Heer. Wrangler: Interactive visual specification of data transformation scripts. In *Proceedings of the sigchi conference on human factors in computing systems*, pp. 3363–3372, 2011. doi: [10.1145/1978942.1979444](https://doi.org/10.1145/1978942.1979444) 1
- [30] H. Kang. Sample size determination and power analysis using the g* power software. *Journal of educational evaluation for health professions*, 18, 2021. doi: [10.3352/jeehp.2021.18.17](https://doi.org/10.3352/jeehp.2021.18.17) 4
- [31] A. Karduni, D. Markant, R. Wesslen, and W. Dou. A bayesian cognition approach for belief updating of correlation judgement through uncertainty visualizations. *IEEE Transactions on Visualization and Computer Graphics*, 27(2):978–988, 2020. doi: [10.1109/TVCG.2020.3029412](https://doi.org/10.1109/TVCG.2020.3029412) 4
- [32] D. H. Kim, V. Setlur, and M. Agrawala. Towards understanding how readers integrate charts and captions: A case study with line charts. In *Proceedings of the 2021 CHI Conference on Human Factors in Computing Systems*, pp. 1–11, 2021. doi: [10.1145/3411764.3445443](https://doi.org/10.1145/3411764.3445443) 3
- [33] J. Klayman. Varieties of confirmation bias. *Psychology of learning and motivation*, 32:385–418, 1995. doi: [10.1016/S0079-7421\(08\)60315-1](https://doi.org/10.1016/S0079-7421(08)60315-1) 3, 5
- [34] J. Klayman and Y.-W. Ha. Confirmation, disconfirmation, and information in hypothesis testing. *Psychological review*, 94(2):211, 1987. doi: [10.1037/0033-295X.94.2.211](https://doi.org/10.1037/0033-295X.94.2.211) 3
- [35] H.-K. Kong, Z. Liu, and K. Karahalios. Frames and slants in titles of visualizations on controversial topics. In *Proceedings of the 2018 CHI Conference on Human Factors in Computing Systems*, pp. 1–12, 2018. doi: [10.1145/3173574.3174012](https://doi.org/10.1145/3173574.3174012) 3
- [36] R. Koonchanok, G. Y. Tawde, G. R. Narayanasamy, S. Walimbe, and K. Reda. Visual belief elicitation reduces the incidence of false discovery. In *Proceedings of the 2023 CHI Conference on Human Factors in Computing Systems*, pp. 1–17, 2023. doi: [10.1145/3544548.3580808](https://doi.org/10.1145/3544548.3580808) 4
- [37] P. Koukaras and C. Tjortjis. Data preprocessing and feature engineering for data mining: Techniques, tools, and best practices. *AI*, 6(10), 2025. doi: [10.3390/ai6100257](https://doi.org/10.3390/ai6100257) 9
- [38] S. Li, T. J. Davidson, C. Xiong Bearfield, and E. Wall. Confirmation bias: The double-edged sword of data facts in visual data communication. In *Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems*, pp. 1–16, 2025. doi: [10.1145/3706598.3713831](https://doi.org/10.1145/3706598.3713831) 3, 5, 6
- [39] M. Lidén. *Confirmation bias in criminal cases*. Oxford University Press, 2023. doi: [10.1093/oso/9780192867643.001.0001](https://doi.org/10.1093/oso/9780192867643.001.0001) 3
- [40] P. C. Mahalanobis. On the generalized distance in statistics. *Sankhyā: The Indian Journal of Statistics, Series A (2008-)*, 80:S1–S7, 2018. doi: [10.1007/s13171-019-00164-5](https://doi.org/10.1007/s13171-019-00164-5) 2
- [41] A. A. Nalcaci, D. Girgin, S. Balki, F. Talay, H. A. Boz, and S. Balcişoy.

- Detection of confirmation and distinction biases in visual analytics systems. In *TrustVis@ EuroVis*, pp. 13–17, 2019. doi: [10.2312/trvis.20191185](https://doi.org/10.2312/trvis.20191185) 3
- [42] R. S. Nickerson. Confirmation bias: A ubiquitous phenomenon in many guises. *Review of general psychology*, 2(2):175–220, 1998. doi: [10.1037/1089-2680.2.2.175](https://doi.org/10.1037/1089-2680.2.2.175) 2, 3
- [43] U. Peters. What is the function of confirmation bias? *Erkenntnis*, 87(3):1351–1376, 2022. doi: [10.1007/s10670-020-00252-1](https://doi.org/10.1007/s10670-020-00252-1) 3
- [44] H. Piringer, W. Berger, and H. Hauser. Quantifying and comparing features in high-dimensional datasets. In *2008 12th International Conference Information Visualisation*, pp. 240–245. IEEE, 2008. doi: [10.1109/IV.2008.17](https://doi.org/10.1109/IV.2008.17) 9
- [45] V. Raman and J. M. Hellerstein. Potter’s wheel: An interactive data cleaning system. In *VLDB*, vol. 1, pp. 381–390, 2001. doi: [10.5555/645927.672045](https://doi.org/10.5555/645927.672045) 1
- [46] W. R. Schumm. Confirmation bias and methodology in social science: an editorial, 2021. doi: [10.1080/01494929.2021.1872859](https://doi.org/10.1080/01494929.2021.1872859) 3
- [47] J. P. Simmons, L. D. Nelson, and U. Simonsohn. False-positive psychology: Undisclosed flexibility in data collection and analysis allows presenting anything as significant. *Psychological science*, 22(11):1359–1366, 2011. doi: [10.1177/0956797611417632](https://doi.org/10.1177/0956797611417632) 2
- [48] J. P. Stevens. Outliers and influential data points in regression analysis. *Psychological bulletin*, 95(2):334, 1984. doi: [10.1037/0033-2909.95.2.334](https://doi.org/10.1037/0033-2909.95.2.334) 2
- [49] J. H. Sullivan, M. Warkentin, and L. Wallace. So many ways for assessing outliers: What really works and does it matter? *Journal of Business Research*, 132:530–543, 2021. doi: [10.1016/j.jbusres.2021.03.066](https://doi.org/10.1016/j.jbusres.2021.03.066) 2
- [50] J. Suschnigg, B. Mutlu, G. Koutroulis, H. Hussain, and T. Schreck. Mandala—visual exploration of anomalies in industrial multivariate time series data. In *Computer Graphics Forum*, p. e70000. Wiley Online Library, 2025. doi: [10.1111/cgf.70000](https://doi.org/10.1111/cgf.70000) 2
- [51] E. R. Tufte and P. R. Graves-Morris. *The visual display of quantitative information*, vol. 2. Graphics press Cheshire, CT, 1983. 2
- [52] E. Wall, L. Blaha, C. Paul, and A. Endert. A formative study of interactive bias metrics in visual analytics using anchoring bias. In *Human-Computer Interaction—INTERACT 2019: 17th IFIP TC 13 International Conference, Paphos, Cyprus, September 2–6, 2019, Proceedings, Part II* 17, pp. 555–575. Springer, 2019. doi: [10.1007/978-3-030-29384-0_34](https://doi.org/10.1007/978-3-030-29384-0_34) 8
- [53] J. Wang, X. Cai, J. Su, Y. Liao, and Y. Wu. What makes a scatterplot hard to comprehend: data size and pattern salience matter. *Journal of Visualization*, 25(1):59–75, 2022. doi: [10.1007/s12650-021-00778-8](https://doi.org/10.1007/s12650-021-00778-8) 9
- [54] L. Wilkinson. Visualizing big data outliers through distributed aggregation. *IEEE transactions on visualization and computer graphics*, 24(1):256–266, 2017. doi: [10.1109/TVCG.2017.2744685](https://doi.org/10.1109/TVCG.2017.2744685) 2, 9
- [55] L. Wilkinson, A. Anand, and R. Grossman. High-dimensional visual analytics: Interactive exploration guided by pairwise views of point distributions. *IEEE Transactions on Visualization and Computer Graphics*, 12(6):1363–1372, 2006. doi: [10.1109/TVCG.2006.94](https://doi.org/10.1109/TVCG.2006.94) 9
- [56] L. Wilkinson and G. Wills. Scagnostics distributions. *Journal of Computational and Graphical Statistics*, 17(2):473–491, 2008. doi: [10.1198/106186008X320465](https://doi.org/10.1198/106186008X320465) 4, 6
- [57] Y. Xiao, S. Ma, A. Oulasvirta, and E. Jun. Priorweaver: Prior elicitation via iterative dataset construction. In *Proceedings of the 2026 CHI Conference on Human Factors in Computing Systems*, number 897, pp. 1–21, 2026. doi: [10.1145/3772318.379034](https://doi.org/10.1145/3772318.379034) 9
- [58] W. Xie, O. Chkrebti, and S. Kurtsek. Visualization and outlier detection for multivariate elastic curve data. *IEEE transactions on visualization and computer graphics*, 26(11):3353–3364, 2019. doi: [10.1109/TVCG.2019.2921541](https://doi.org/10.1109/TVCG.2019.2921541) 2
- [59] C. Xiong, C. Stokes, Y.-S. Kim, and S. Franconeri. Seeing what you believe or believing what you see? belief biases correlation estimation. *IEEE Transactions on Visualization and Computer Graphics*, 2022. to appear. doi: [10.1109/TVCG.2022.3209405](https://doi.org/10.1109/TVCG.2022.3209405) 1, 6
- [60] I. Yanai and M. J. Lercher. Does the data do the talking? confirmation bias in the interpretation of data visualizations. 2021. doi: [10.31234/osf.io/q8ckd](https://doi.org/10.31234/osf.io/q8ckd) 3
- [61] J. S. Yi, Y. ah Kang, J. Stasko, and J. A. Jacko. Toward a deeper understanding of the role of interaction in information visualization. *IEEE transactions on visualization and computer graphics*, 13(6):1224–1231, 2007. doi: [10.1109/TVCG.2007.70515](https://doi.org/10.1109/TVCG.2007.70515) 4