

Towards Understanding How Readers Integrate Charts and Captions: A Case Study with Line Charts

Dae Hyun Kim
dhkim16@cs.stanford.edu
Stanford University/Tableau Research
Stanford/Palo Alto, California

Vidya Setlur
vsetlur@tableau.com
Tableau Research
Palo Alto, California

Maneesh Agrawala
maneesh@cs.stanford.edu
Stanford University
Stanford, California

ABSTRACT

Charts often contain visually prominent features that draw attention to aspects of the data and include text captions that emphasize aspects of the data. Through a crowdsourced study, we explore how readers gather takeaways when considering charts and captions together. We first ask participants to mark visually prominent regions in a set of line charts. We then generate text captions based on the prominent features and ask participants to report their takeaways after observing chart-caption pairs. We find that when both the chart and caption describe a high-prominence feature, readers treat the doubly emphasized high-prominence feature as the takeaway; when the caption describes a low-prominence chart feature, readers rely on the chart and report a higher-prominence feature as the takeaway. We also find that external information that provides context, helps further convey the caption's message to the reader. We use these findings to provide guidelines for authoring effective chart-caption pairs.

CCS CONCEPTS

• **Human-centered computing** → **Empirical studies in visualization**.

KEYWORDS

Captions; line charts; visually prominent features; takeaways.

ACM Reference Format:

Dae Hyun Kim, Vidya Setlur, and Maneesh Agrawala. 2021. Towards Understanding How Readers Integrate Charts and Captions: A Case Study with Line Charts. In *CHI Conference on Human Factors in Computing Systems (CHI '21)*, May 8–13, 2021, Yokohama, Japan. ACM, New York, NY, USA, 11 pages. <https://doi.org/10.1145/3411764.3445443>

1 INTRODUCTION

Charts provide graphical representations of data that can draw a reader's attention to various visual features such as outliers and trends. Readers are initially drawn towards the most *visually salient* components in the chart such as the chart title and the labels [29]. However, they eventually apply their cognitive processes to extract meaning from the most *prominent* chart features [5, 41]. Consider



the line chart at the beginning of this article. What do you think are the main visual features of the chart and what are its key takeaways?

Such charts are often accompanied by text captions that emphasize specific aspects of the data as chosen by the chart author. In some cases, the data emphasized in the caption corresponds to the most visually prominent features of the chart and in other cases it does not. Prior studies have shown that charts with captions can improve both recall and comprehension of some aspects of the underlying information, compared to seeing the chart or the caption text alone [3, 15, 24, 32]. But far less is known about how readers integrate information between charts and captions, especially when the data emphasized by the visually prominent features of the chart differs from the data that is emphasized in the caption.

Consider the visually prominent features in our initial line chart and then consider each of the following caption possibilities one at a time. How do your takeaways change with each one?

- (1) The chart shows the 30-year fixed mortgage rate between 1970 and 2018.
- (2) The 30-year fixed mortgage rate increased slightly from 1997 to 1999.
- (3) The 30-year fixed mortgage rate reached its peak of 18.45% in 1981.
- (4) The 30-year fixed mortgage rate reached its peak of 18.45% in 1981 due to runaway inflation.

The first caption simply describes the dimensions graphed in the chart and only provides redundant information that could be read from the axis labels. Automated caption generation tools often create such *basic* descriptive captions [30, 40]. The next three captions each emphasizes aspects of the data corresponding to a visual feature of the chart (i.e., upward trend, peak) by explicitly mentioning the corresponding data point or trend. However, the second caption emphasizes a feature of low visual prominence – a relatively local and small rise in the chart between 1997 and 1999. The third caption describes the most visually prominent feature of the chart – the tallest peak that occurs in 1981. The fourth caption also describes this most visually prominent feature, but adds external information that is not present in the chart and provides context for the data.

In this paper, we examine two main hypotheses - (1) When a caption emphasizes more visually prominent features of the chart,

Permission to make digital or hard copies of part or all of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for third-party components of this work must be honored. For all other uses, contact the owner/author(s).

CHI '21, May 8–13, 2021, Yokohama, Japan

© 2021 Copyright held by the owner/author(s).

ACM ISBN 978-1-4503-8096-6/21/05.

<https://doi.org/10.1145/3411764.3445443>

people are more likely to treat those features as the takeaway; even when a caption emphasizes a less visually prominent feature, people are still more likely to treat a more visually prominent feature in the chart as the takeaway. (2) When a caption contains external information for context, the information serves to further emphasize the feature described in the caption and readers are therefore more likely to treat that feature as the takeaway.

We considered univariate line charts for our work because they are among the most common basic charts and are easily parameterizable, making them useful for the initial exploration of our hypotheses. We synthesized 27 single charts with carefully chosen parameters and collected 16 real-world single line charts to confirm the generalizability of our findings. We ran a data collection activity on the 43 single-line charts, where we asked 219 participants to mark visually prominent regions on the line charts. We generated text captions for the ranked set of prominent features using templates to control variations in natural language. Finally, we conducted a crowdsourced study with a new set of 2168 participants to report their takeaways after seeing the chart-caption pairs.

Our findings from the study support both of our hypotheses. Referring back to our initial line chart, when the caption mentions the most prominent feature as in the third caption (i.e., the peak in 1981), readers will probably take away information from that feature. When the caption mentions a less prominent feature as in the second caption (i.e., the increase from 1997 to 1999), there is a mismatch in the message between the chart and the caption. Readers will have a strong tendency to go with the message conveyed in the chart and take away information about the peak value. Finally, the external information about the peak value present in the fourth caption will reinforce the message in the caption and the readers will more likely take away information about the peak.

These findings help better understand the relationship between charts and their captions when conveying information about certain aspects of the data to the reader. Based on these studies, we provide guidelines for authoring charts and captions together in order to emphasize the author's intended takeaways. Visualization authors can more effectively convey their message to readers by ensuring that both charts and captions emphasize the same set of features. Specifically, authors could make visual features that are related to their key message, more prominent through visual cues (e.g., highlighting or zooming into a focus area, adding annotations) [10, 26] or include external information in the caption to further emphasize the feature described in the caption. Often, an alternative chart representation may be more conducive to making certain visual features more prominent.

2 RELATED WORK

Our work is related to two lines of research: (1) Cognitive Understanding of Charts and (2) Caption Generation Tools.

2.1 Cognitive Understanding of Charts

The prevalence of text with visuals has led researchers to explore how readers specifically understand information in figures with accompanying text in several domains. Li et al. [25] conducted studies to demonstrate that figures with text can convey essential information and better aid understanding than just text alone for scientific

publications in a biomedical domain. Odell et al. [34] demonstrated that having text that accurately describes important findings in medical diagnostic images can increase physicians' speed and accuracy on Bayesian reasoning tasks while making life-critical judgments for patients. Xiong et al. [47] showed that background knowledge can affect viewers' visual perception of data as they tend to see the pattern in the data corresponding to the background knowledge as more visually salient. Kong et al. [20] explored the impact of titles on visualization interpretation with different degrees of misalignment between a visualization and its title. A title contains a miscued slant when the visualization emphasizes one side of the story through visual cues but the title's message addresses the other (less emphasized) side of the story. Titles have a contradictory slant where the information conveyed in the title is not presented at all in the visualization. They observe that even though the title of a visualization may not be recalled, the title can still measurably impact the remembered contents of a chart. Specifically, titles with a contradictory slant trigger more people to identify bias compared to titles with a miscued slant, while visualizations are perceived as impartial by the majority of viewers [21]. Elzer et al. [11] conducted a study to better understand the extent to which captions contribute to recognizing the intended message of an information graphic for sight-impaired users. They find that the caption strengthens the intended message of the graphic. Carberry et al. [4] showed that the communicative goals of infographics in digital libraries are often not repeated in the text of the articles. Their work looked into how information in the graphics could be better utilized for summarizing a document by employing a Bayesian network.

However, this previous research has not explored the relationship between charts and their captions with respect to how they work together to emphasize certain aspects of the data to the reader.

2.2 Caption Generation Tools

A number of visual analysis tools help users design charts and captions from an input data table [8, 9, 16, 42, 46]. These captions generally only describe the data attributes and visual encodings that are in play in the charts and do not highlight key takeaways. Nevertheless, authors often include text with a chart to help emphasize an intended message to their audience. PostGraphe [13] generated reports integrating graphics and text from a list of user-defined intentions about the underlying data such as the comparison of variables. SAGE [31] used natural language generation techniques to produce explanatory captions for information graphics. The system generates captions based on the structural and spatial relations of the graphical objects and their properties along with explanations describing the perceptual complexity of the data attributes in the graphics. SumTime [48] used pattern recognition techniques to generate textual summaries of time-series data. The iGRAPH-Lite system [14] made information in a graphic accessible to blind users by using templates to provide textual summaries of what the graphic looks like. The summaries however, do not focus on the higher-level takeaway conveyed by the graphic. Chen et al. [7] produced natural language descriptions for figures by identifying relations among labels present in the figures.

Other work has explored natural language generation techniques for assembling multiple caption units together to form captions [37].

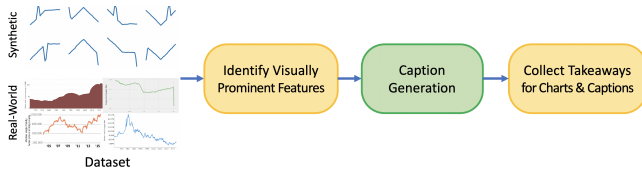


Figure 1: Our study pipeline. The inputs to the study are 27 synthetic and 16 real-world charts. Yellow boxes represent steps where we employed crowdsourcing. The green box indicates that the step did not involve crowdsourcing.

Deep learning techniques based on neural networks automate caption generation tasks for news images [6]. Elzer et al. [12] identified communicative signals that represent the intent of messages portrayed in basic bar charts by applying a Bayesian network methodology for reasoning about these signals and generating captions. Liu et al. [27, 28, 44] explored the integration of text analytics algorithms with interactive visualization tools to help users understand and interpret the summarization results. Contextifier [17] automatically annotated visualizations of stock behavior with news article headlines taking into account visual salience contextual relevance, and key events from the articles. Kim et al. [18] introduced an automatic chart question answering pipeline that generates visual explanations that refer to visual features of charts using a template-based natural language generation approach. Voder [38] generated data facts for visualizations with embellishments to help users interpret visualizations and communicate their findings. While an evaluation of that system suggested that interactive data facts aided users in interpreting visualizations, the paper did not specifically explore the interplay between data facts and the visualizations and their effects on the readers’ takeaways.

These systems focus on helping authors with auto-generated text that can be associated with graphics; however, the work does not evaluate what information readers gather from the generated captions with their corresponding graphics. Our paper specifically explores how similarities and differences between what is visually emphasized in a line chart and textually emphasized in its caption, can affect what readers take away from the information when presented together. Future directions from our work could extend the functionality of chart authoring tools by providing automatic suggestions for captions as well as for chart presentation to help the reader take away information that is consistently emphasized by both the chart and caption.

3 STUDY

We conducted a crowdsourced study to understand how captions describing features of varying prominence levels and the effect of including or not including external information for context, interacts with the chart in forming the readers’ takeaways. Through an initial data collection activity, we asked participants to identify features in the line charts that they thought were visually prominent. We generated captions corresponding to those marked features of various levels of prominence. We then ran a study asking a new set of participants to type their takeaways after viewing a chart and caption pair. Figure 1 shows the study pipeline.

3.1 Datasets

We ran the study on two different datasets - (1) synthetically generated line charts that we designed to ensure good coverage of a variety of visual features that occur in line charts and (2) line charts gathered from real-world sources to serve as a more ecologically valid setting for our study.

Synthetic Charts. We generated a set of synthetic line charts with common visual features (i.e., trends, extrema, and inflection points) while maintaining realistic global shapes. To keep the overall design space tractable, we limited global shapes to include at most two trends (i.e., up, down, and flat) and added at most one perturbation to induce features (e.g. inflection points) in either the positive or negative direction, resulting in a total of 27 data shapes (Figure 2). To provide context to the charts, we labeled the x-axis with time unit values implying that the chart represents a time series. Specifically, we selected the start and end of the x-axis from the set of years {1900, 1910, 1920,..., 2020}. To label the y-axis, we chose a domain for the y-axis and its value range from the MassVis dataset [2].

Real-world Charts. To build a more ecologically representative dataset of line charts with various shapes, styles, and domains, we collected 16 charts (Figure 3) from sources such as The Washington Post [43], Pew Research [36], Wikipedia [45], and Tableau Public [39]. Because our study focuses on prominence arising from intrinsic features in line charts, we removed all graphical elements that could potentially affect the prominence of the features in the charts (e.g., text annotations, highlighting, and background shading). In addition, we removed all text except for the axis labels (e.g. chart titles) so that the captions serve as the primary source of text provided with the chart. We added axis labels to those charts without labels to ensure readability.

3.2 Identify Visually Prominent Features

To identify the most visually prominent features in our dataset, we recruited at least five workers from Amazon Mechanical Turk [1] for each line chart and asked them to draw rectangular bounding boxes around the top three most prominent features in the chart. We also asked them to briefly describe each marked feature in their own words so that we could differentiate between trend and slope features versus peak, inflection, and other point features.

In each trial of the data collection, we presented one of the 43 line charts. Because we were seeking subjective responses, each participant completed only one trial to avoid biases that might arise from repeated exposure to the task. Participation was limited to English speakers in the U.S. with at least a 98% acceptance rate and 5000 approved tasks. We paid a rate equivalent to \$2 / 10 mins.

We asked a total of 219 participants (average of 5.09 per chart) to label the top three features for a total of 657 prominence boxes. We then aggregated all of the feature bounding boxes provided by first projecting each box onto the x-axis, to form a 1D interval (Figure 4 upper left). We weighted each interval inversely proportional to the ranking provided by the participant. Specifically, the top ranked feature bounding box for each participant was assigned a weight of 3, while the 3rd ranked feature was assigned a weight of 1. We noticed that bounding boxes corresponding to the same features were pretty consistent in the central regions although the

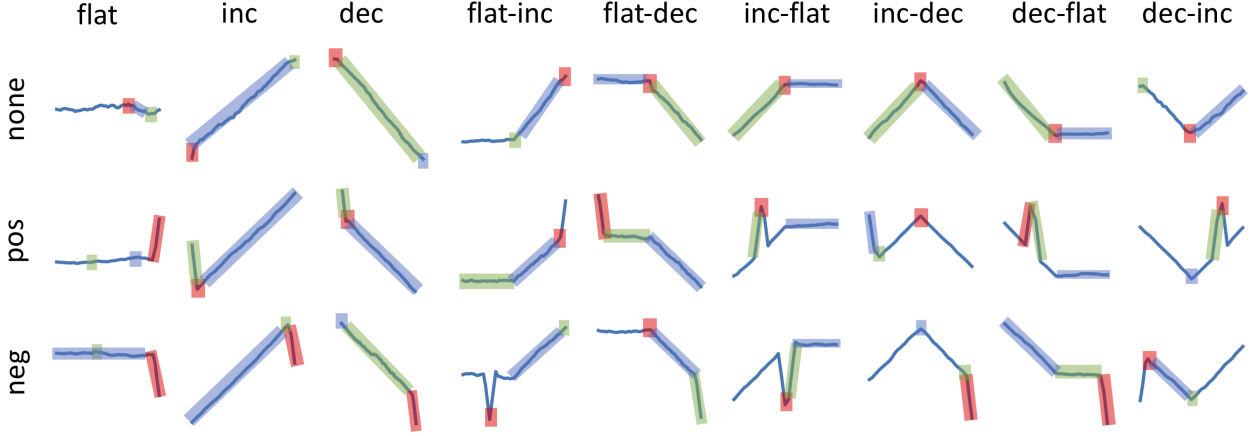


Figure 2: The 27 data shapes generated for the study and their top three prominent features. Columns represent the nine possible global shapes and rows represent the three possible local outlier types. Here, ‘flat’, ‘inc’, and ‘dec’ denote flat, increasing, and decreasing trends respectively. ‘none’, ‘neg’, and ‘pos’ denote none, negative, and positive outlier types respectively. Red, green, and blue regions indicate the top three prominent features in order.

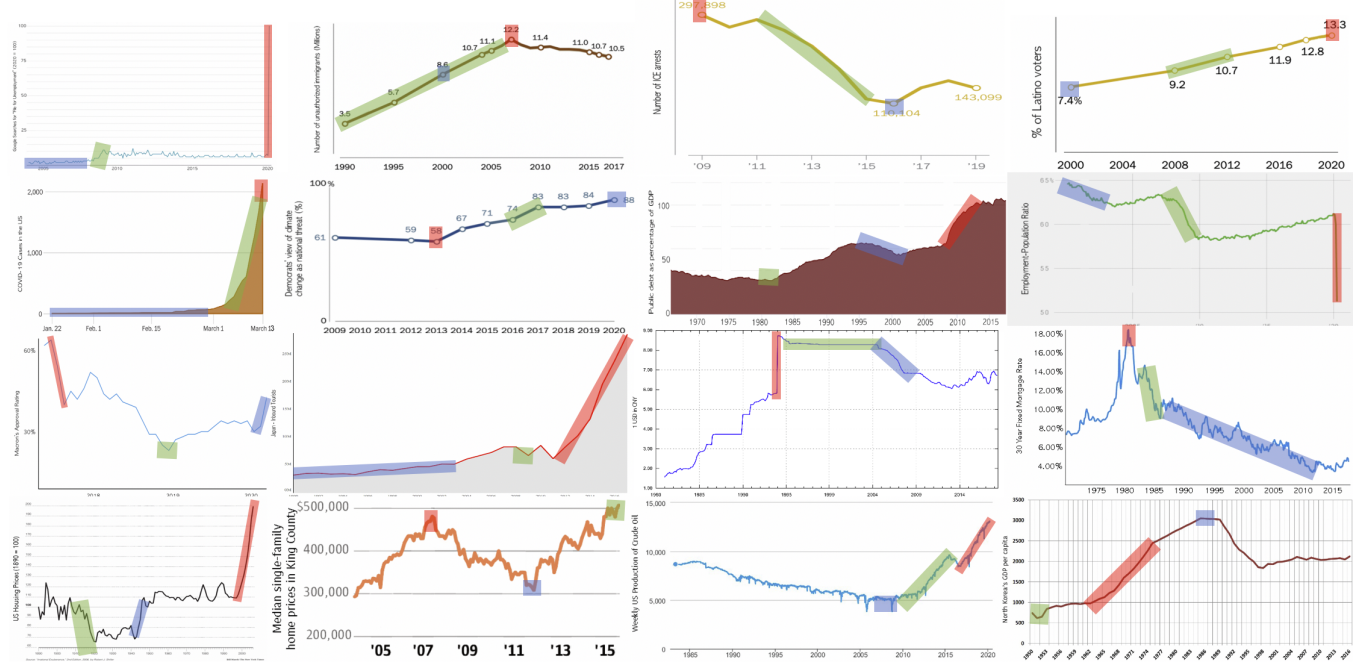


Figure 3: The 16 real-world charts. Red, green, and blue regions indicate the top three prominent features in order.

exact boundary drawn by the participants varied. In order to boost the signal in the central regions while suppressing the noise in the boundary regions, we multiplied the weight assigned to each interval by a Gaussian factor centered at the interval and with standard deviation set to half the width of the interval. Summing all of the Gaussian weighted intervals, we obtained a *prominence curve* (Figure 4 bottom left). However, a region defined by a local maximum of the curve may not have an obvious one-to-one mapping with a feature in the chart because it roughly indicates a high prominence region instead of pinpointing a specific visual feature. We considered all the bounding boxes containing the region along

with the participants’ text descriptions of the features to associate the local maximum to a certain feature. We iterated this process for the region around the top three local maximum to identify three prominent features. Results of the algorithm for the charts in our dataset are shown in Figures 2 and 3.

3.3 Caption Generation

To carefully control the language used in the captions and keep the number of conditions manageable, we generated captions using templates that only vary the feature mentioned and whether

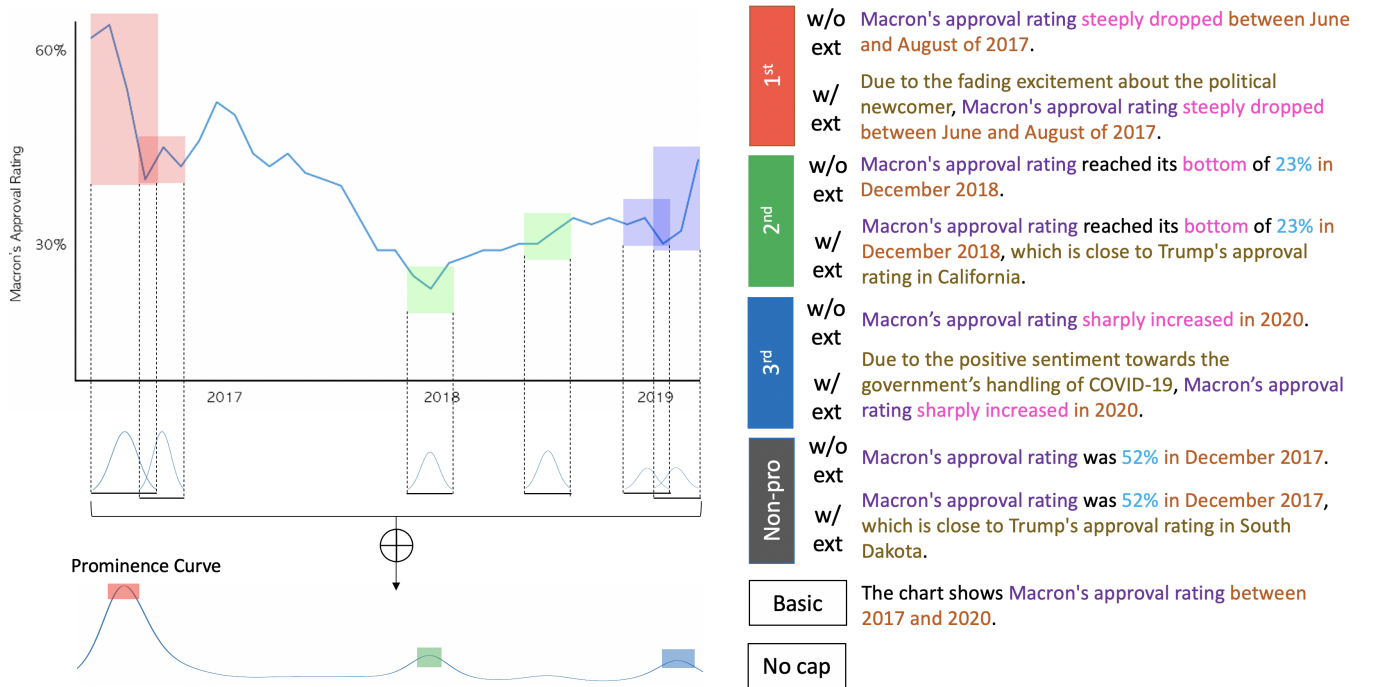


Figure 4: The line on the bottom left shows the prominence curve for the line chart above. From this curve, we obtain the most prominent (red), the second most prominent (green), and the third most prominent (blue) features in the chart. The 10 caption variants (one of them being a no-caption variant) generated based on these prominent features, are shown on the right. The text colors indicate the types of fill-in values based on the caption templates; **purple** for dimensions, **fuchsia** for the feature description, **blue** for data values, and **brown** for the time period.

external information is introduced. Using the templates, we produced the following caption variants: (1) two captions (one with and one without external information) for each of the top three visually prominent features identified earlier, (2) two captions (one with and one without external information) describing a minimally prominent feature that is neither an extremum nor an inflection point, and (3) a basic caption that simply describes the domain represented in the chart without describing a particular feature.

We generated 10 caption variants (including the no caption variant in which we presented a chart without caption) for each of the 43 charts, providing a total of 430 chart-caption pairs. We manually generated all the captions rather than using the original captions for the real-world charts to control for word use and grammatical structure. For real-world charts, we searched for information from the document that they originally appeared in, to extract information not present in the charts. In particular, we looked for information about potential reasons for trends or change (e.g. the external information included in the caption about the most prominent feature in Figure 4) or comparisons with a similar entity (e.g. comparison between Macron’s approval rating with Trump’s approval rating in the second most prominent feature in Figure 4). For synthetically generated charts and real-world charts that were not accompanied with additional information about their features, we referenced Wikipedia [45] articles to create a plausible context.

We employed simple language templates for caption generation to minimize the effects of linguistic variation (Table 1). The captions generated with the templates were allowed to vary in the features

Feature	Template
Extremum	[dimension] reached its [extrema-word] of [value] in [time-period].
Trend	[dimension] [slope-word] in / between [time-period].
Inflection	[dimension] started [slope-word] in [time-period].
Point	[dimension] was [value] in [time-period].

Table 1: Examples of templates we employed for generating captions about specific features. The text colors indicate the types of fill-in values based on the caption templates; **purple** for dimensions, **fuchsia** for feature descriptions, **blue** for data values, and **brown** for time periods. Examples of filled in captions are in Figure 4 (right).

they described in the charts. To make the descriptions of the features appear natural, we included words the participants used to describe the features during the prominent feature collection phase. Because the participants usually described each of the features using a noun occasionally with an adjective modifier (e.g. “sharp increase”), we manually lemmatized the words and modified the forms to correctly fit into our template (e.g. “sharply increased” in the caption about the third most prominent figure in Figure 4).

3.4 Collect Takeaways for Charts & Captions

3.4.1 Design. We ran a between-subjects design study for collecting takeaways for charts and their captions. For each of the 43 charts, we presented one of the ten variants (including the no caption variant) (examples in Figure 4):

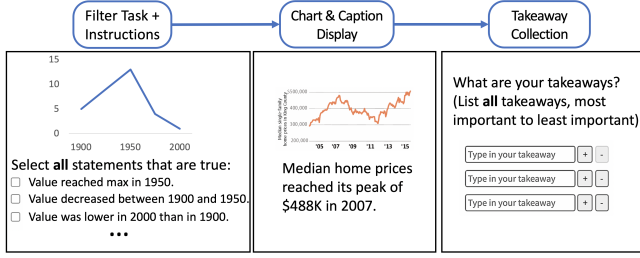


Figure 5: The procedure for collecting takeaways for chart-caption pairs. The images show simplified versions of the screen that the participants saw during each step.

- (1) [1st w/o ext] Caption for most prominent feature, no external info.
- (2) [1st w/ ext] Caption for most prominent feature, has external info.
- (3) [2nd w/o ext] Caption for 2nd most prominent feature, no external info.
- (4) [2nd w/ ext] Caption for 2nd most prominent feature, has external info.
- (5) [3rd w/o ext] Caption for 3rd most prominent feature, no external info.
- (6) [3rd w/ ext] Caption for 3rd most prominent feature, has external info.
- (7) [non-pro w/o ext] Caption for non-prominent feature, no external info.
- (8) [non-pro w/ ext] Caption for non-prominent feature, has external info.
- (9) [basic] Caption about domain represented in the chart and x-range
- (10) [no cap] No caption

3.4.2 Procedure. The study began with a screening test to ensure that the participant had a basic understanding of line charts and could read values and encodings, extract extrema and trends, and compare values (Figure 5 first step). Only participants who passed this test were allowed to continue with the study. After they read the instructions, the participants were presented with a chart and a caption underneath the chart, similar to most charts in the real world (unless it is the no-caption variant) (Figure 5 second step). We did not impose a time constraint on the amount of time spent looking at the chart and the caption to allow participants sufficient time to read and digest the information at their own pace, like document reading in the real world. On the next screen for collecting takeaways, the chart and the caption were removed to constrain readers to provide the takeaways based on memory instead of simply re-reading from the chart and the caption. The participants were asked to list as many text takeaways as they could in the order of importance (Figure 5 third step). Finally, using a 5-point Likert scale, we asked how much they relied on the chart and caption individually when determining their takeaways.

We asked each participant to provide takeaways for exactly one chart-caption pair to prevent potential biases from already having read a different caption about a chart. From 2168 participants (average of 5.04 per chart-caption pair), we collected a total of 4953 takeaways (average of 2.28 per participant).

3.4.3 Labeling Takeaways. In order to analyze the takeaways, we manually labeled each takeaway with the corresponding chart feature described. Since participants often described multiple chart features in a single takeaway, we first split each takeaway into separate takeaways for each visual feature mentioned. At the end of this process, we identified on average 1.31 features per takeaway. If the referenced feature was one of three most prominent features or

the non-prominent feature we identified during caption generation, we labeled the takeaway with the corresponding feature, otherwise we labeled the takeaway as referring to an *other* feature. If the takeaway did not refer to any specific feature in the chart, we labeled the takeaway as a *non-feature*. Examples of *non-feature* takeaways include an extrapolation such as “The value will continue to rise after 2020” or a judgment such as “I should buy gold” when looking at a chart showing the price of gold over time. One of the authors labeled the features and discussed any confusing cases with the other authors to converge on the final label.

4 RESULTS

The primary goal of our study is to understand what readers take away when charts and captions are presented together and how the emphasis on different prominent features and presence of external information affects the takeaways. We analyze our results with respect to two hypotheses:

[H1] When captions emphasize more visually prominent features of the chart, people are more likely to treat the features as the takeaway; when a caption emphasizes a less visually prominent feature, people are less likely to treat that feature as the takeaway and more likely to treat a more visually prominent feature in the chart as the takeaway.

[H2] When captions contain external information for context, the external information serves to further emphasize the feature presented in the caption and people are therefore more likely to treat that feature as the takeaway, compared to when the caption does not contain external information.

Assessing H1. To evaluate **H1**, we examine how varying the prominence of a visual feature mentioned in a caption (independent variable), affects the visual feature mentioned in the takeaways (dependent variable). Figure 6 summarizes the study results for the synthetic charts (top row) and the real-world charts (bottom row).

In general, these results suggest that when a caption mentions visual features of differing prominence levels, the takeaways also differ. Omnibus Pearson’s chi-squared tests confirm a significant difference between the bar charts for the 5 different caption conditions in both the synthetic ($\chi^2(20) = 202.211, p < 0.001$) and real world ($\chi^2(20) = 207.573, p < 0.001$) datasets. These results also suggest that when the caption mentions a specific feature, the takeaways also tend to mention that feature, when compared to the baseline ‘no-caption’ condition.

Figures 7a and 7b collect the percentage of takeaways that mention the same feature as in the caption for the synthetic and the real-world datasets respectively (left darker bars) and compare them with the percentages corresponding to the no-caption case (lighter-hued bars on the right). We see that captions do play a role in forming takeaways and the takeaway is thus more likely to mention that feature (i.e., each darker bar in Figures 7a and 7b is usually longer than the corresponding lighter-hued bar to its right). Planned pairwise Z-tests with Bonferroni correction are shown in Table 2. Block 1 shows that the differences between the corresponding color bars are significant for the second most prominent, third most prominent, and non-prominent features. For the most prominent feature, we find that while a higher proportion of people

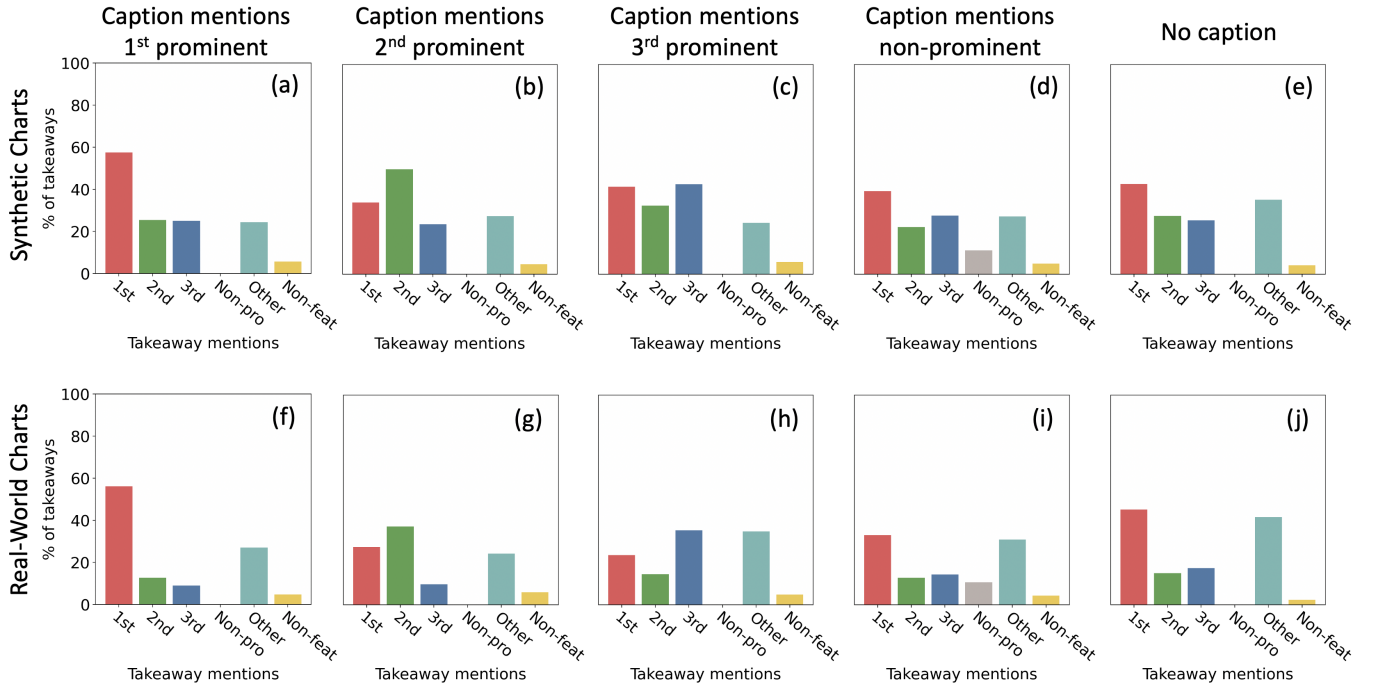


Figure 6: Study results. Each column shows bar charts for each prominence level mentioned in the caption (i.e., the leftmost bar chart is for captions mentioning the 1st ranked visual feature, the next bar chart is for captions mentioning the 2nd ranked visual feature, while the rightmost bar chart is for the no-caption condition). Within a bar chart, each bar represents the percentage of takeaways mentioning the visual feature at that prominence level. For example, the leftmost bar in each bar chart represents the percentage of total takeaways that mention the top ranked takeaway. Each bar chart also reports the percentage of *Other* features and *Non-features* that were mentioned in the takeaways. These charts aggregate data for captions with and without external information. The percentages do not sum to 100% as some takeaways mention multiple takeaways.

mentioned the most prominent feature in their takeaways when the caption mentions it, the difference is only significant for the synthetic charts. We believe that this is possibly because people already include the most prominent features in their takeaways in the no-caption condition and the difference hence is not significant.

While we confirmed that both the chart and caption play a role as to what the reader takes away from them, the key question is how the chart and the caption interact with each other – Do they have a synergistic effect when they emphasize the same feature? Which one wins over when they emphasize different features? Referring to Figure 6, we see the synergistic effect of the double-emphasis from the chart and caption when they emphasize the same feature (Figures 6a and 6f). In particular, the participants took away from the most prominent feature significantly more often than from any other feature in the chart (Table 2 Block 3). When the caption diverged from the chart and described a feature that was not prominent, the participants relied more on the chart and took away from the most prominent feature significantly more than the feature described in the caption (Table 2 Block 4, rows 3 and 6; Figures 6d and 6i). When the caption did not diverge as much and described the second or the third most prominent feature, the takeaways mentioned the feature described in the caption more than the most prominent feature (Table 2 Block 4, rows 1, 2, 4, and 5; Figures 6b, 6c, 6g, and 6h). However, the difference was smaller than the difference between the ratio of people who took away from

the most prominent feature and the ratio of people who took away from any of the other features. We believe this result may be due to the fact that the charts still had more influence on the readers than the captions as the second and the third most prominent feature are still among the top prominent features and are among the features emphasized by the chart.

We observe from Figure 7 that the chart also plays an important role in what people take away – when a caption mentions a higher-prominent feature, the takeaways more consistently mention that feature. Specifically, we see that the bars for the higher-prominence features are taller than the bars for the lower-prominence features, indicating an increase in the effectiveness of chart in reinforcing the message in the caption. Planned pairwise Z-tests with Bonferroni correction between each subsequent pair of bars (red bar vs. green bar, green bar vs. blue bar, blue bar vs. gray bar) (Table 2 Block 2) find that the red bar vs. green bar is significant for real-world charts and the blue bar vs. gray bar is significant both synthetic and real-world charts, whereas the green bar vs. blue bar difference is not significant. We believe that the visual prominence levels for some of the top-ranked features are similar in several charts (i.e., the difference in prominence between the 1st and 2nd ranked features is small) in our dataset and this results in a smaller difference between them, although the trend is in the right direction.

Table 3 shows average and standard deviation of how much the participants reported to have relied on the chart and the caption

Source	Caption-Takeaway 1		Caption-Takeaway 2		Z	p
	Caption	Takeaway	Caption	Takeaway		
Block 1. Takeaways mentioning feature in caption vs. without caption						
Synthetic	1st	1st	no cap	1st	2.846	0.002*
	2nd	2nd	no cap	2nd	4.641	< 0.001*
	3rd	3rd	no cap	3rd	3.643	0.001*
	non-pro	non-pro	no cap	non-pro	6.195	< 0.001*
Real-world	1st	1st	no cap	1st	1.660	0.049
	2nd	2nd	no cap	2nd	4.225	< 0.001*
	3rd	3rd	no cap	3rd	3.347	< 0.001*
	non-pro	non-pro	no cap	non-pro	4.732	< 0.001*
Block 2. Between takeaways mentioning feature in caption						
Synthetic	1st	1st	2nd	2nd	1.782	0.037
	2nd	2nd	3rd	3rd	0.705	0.044
	3rd	3rd	non-pro	non-pro	8.989	< 0.001*
Real-world	1st	1st	2nd	2nd	3.708	< 0.001*
	2nd	2nd	3rd	3rd	0.363	0.358
	3rd	3rd	non-pro	non-pro	5.940	< 0.001*
Block 3. When caption = 1st: takeaway = 1st vs. takeaway $\neq$ 1st						
Synthetic	1st	1st	1st	2nd	8.168	< 0.001*
	1st	1st	1st	3rd	8.275	< 0.001*
	1st	1st	1st	non-pro	19.463	< 0.001*
Real-world	1st	1st	1st	2nd	9.981	< 0.001*
	1st	1st	1st	3rd	11.301	< 0.001*
	1st	1st	1st	non-pro	11.536	< 0.001*
Block 4. When caption $\neq$ 1st: takeaway = 1st vs. takeaway = caption						
Synthetic	2nd	2nd	2nd	1st	3.829	< 0.001*
	3rd	3rd	3rd	1st	0.258	0.398
	non-pro	1st	non-pro	non-pro	8.342	< 0.001*
Real-world	2nd	2nd	2nd	1st	2.010	0.022
	3rd	3rd	3rd	1st	2.521	0.006*
	non-pro	1st	non-pro	non-pro	5.454	< 0.001*

Table 2: Pairwise Z-test results of comparisons between various ratios of takeaways that mention a certain feature (third, fifth columns) when provided a caption describing a certain feature (second, fourth columns). The tests were one-sided with the alternative hypothesis that the ratio of takeaways for ‘Caption-Takeaway 1’ is greater than the ratio of takeaways for ‘Caption-Takeaway 2’. Asterisks indicate significance with Bonferroni correction.

respectively on a 5-point Likert scale. The results in Table 3 Block 1 suggest that the participants drew information from both the chart and the caption when determining their takeaways, although they consistently relied on the chart more than the caption. These results potentially shed light on why participants took away more often from the chart than the caption when they began to diverge – they relied more on the chart than the caption. The results further suggest that the participants’ tendency to rely on the charts grew while their tendency to rely on the captions declined as the prominence of the feature described in the caption decreased (Table 3 Block 2). We found a significant drop in the self-reported reliance on the caption when the caption described a non-prominent feature compared to when it described the third-most prominent feature (synthetic: Mann-Whitney $U = 28941$, $p < 0.001$; real-world: Mann-Whitney $U = 9666$, $p < 0.001$) whereas the increase in the reported reliance on the chart when the caption described a non-prominent feature was only significant with the synthetic charts (Mann-Whitney $U = 32844.5$, $p < 0.001$). Although the general trend is in the right direction, we did not find significant differences in the reliance scores when the caption mentioned one of the top three prominent features. This may be because the difference in prominence is not as great among these features as it is with the non-prominent feature. These results are in line with our findings from the takeaways; we

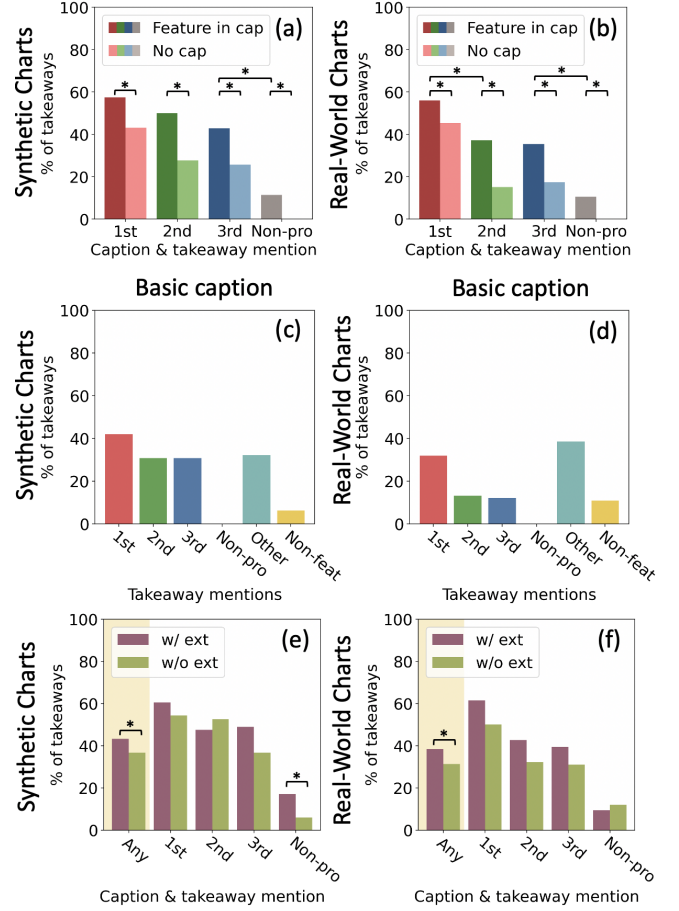


Figure 7: (Top row) Comparison of percentages of takeaways that mention the same feature as the caption for the synthetic (a) and real-world (b) datasets (i.e., darker bars on the left correspond to the red bar from Figure 6a, the green bar from 6b, the blue bar from 6c, and the grey bar from 6d), and percentages of takeaways that mention the feature in the no caption condition (i.e., the right lighter-hued bars in the chart correspond to the bars from Figure 6e). (Middle row) Percentage of takeaways mentioning the visual features at each prominence level when presented with the basic caption. (Bottom row) Dividing the left bars in charts (top row)a and (top row)b based on whether the caption contains external information (purple bars) or does not (olive bars). The leftmost Any bars show aggregates over all prominence levels. Asterisks indicate significant difference.

find that when the chart contains a high-prominence visual feature, but the caption emphasizes a low-prominence feature, participants relied more on the chart and less on the caption.

Considering all these results together suggests that we can accept our hypothesis **H1** – readers take away from the highly prominent features when the chart and caption both emphasize the same feature and that their inclination to rely more on the most prominent feature instead of the feature described in the caption becomes greater when the caption describes a less prominent feature.

Source	Caption Type	Reported Reliance	
		Chart	Caption
Block 1. Overall			
Synthetic	all	4.675 ± 0.670	2.624 ± 1.609
Real-world	all	4.536 ± 0.784	2.779 ± 1.679
Block 2. Prominence			
Synthetic	1st	4.590 ± 0.711	3.249 ± 1.327
	2nd	4.567 ± 0.814	3.082 ± 1.433
	3rd	4.567 ± 0.726	3.059 ± 1.408
	non-pro	4.775 ± 0.549	2.447 ± 1.429
	basic	4.850 ± 0.377	2.593 ± 1.320
Real-world	1st	4.494 ± 0.838	3.405 ± 1.481
	2nd	4.462 ± 0.890	3.165 ± 1.359
	3rd	4.503 ± 0.805	3.236 ± 1.354
	non-pro	4.595 ± 0.718	2.680 ± 1.545
	basic	4.628 ± 0.601	2.718 ± 1.568
Block 3. External Information			
Synthetic	w/o ext	4.679 ± 0.688	2.798 ± 1.402
	w/ ext	4.573 ± 0.728	3.110 ± 1.448
Real-world	w/o ext	4.606 ± 0.741	3.061 ± 1.481
	w/ ext	4.424 ± 0.875	3.194 ± 1.439

Table 3: The reported reliance on the chart and the caption respectively on 5-point Likert scales. Block 1 shows the reported reliance across all the captions. Block 2 shows the reported reliance depending on the prominence of the feature described in the chart and Block 3 shows the reported reliance depending on the inclusion of external information. The values are reported in the form of $\mu \pm \sigma$.

H1 Additional Results. We also collected takeaways for charts with *basic* captions that describe the axes of the chart. (Figure 7 - middle row). We find that the percentage of takeaways for each of the features is similar to that of the no-caption condition. In fact, Pearson’s chi-square test finds no significant difference between the takeaway histograms of the basic caption and the no-caption conditions (synthetic: $\chi^2(4) = 1.564$, $p = 0.815$; real-world: $\chi^2(4) = 7.168$, $p = 0.127$). While automated captioning tools [30, 40] generate captions corresponding to our basic captions, we were unable to find evidence that these captions affect what people take away. Such captions may help readers with accessibility needs; however, we believe further exploration will help future systems determine appropriate uses for such captions.

Assessing H2. To evaluate H2, we examine whether including external content information in the caption makes it more likely for readers to take away the feature mentioned in the caption. We find that people are significantly more likely to mention the feature described in the caption when it includes external information than when it does not (Figures 7e and Figures 7f *Any* bars). A pairwise Z-test finds significant difference between these ratios (synthetic: $Z = 2.273$, $p = 0.011$; real-world: $Z = 2.032$, $p = 0.021$). In addition, the reported reliance on the chart and the captions shifted towards the captions with external information, which is in-line with our findings (Figure 3 Block 3). Specifically, the reported reliance on the chart was significantly lower with external information (synthetic: Mann-Whitney $U = 137318$, $p < 0.001$; real-world: Mann-Whitney $U = 45292$, $p = 0.001$); the reported reliance on the caption was

higher with external information, but the difference was only significant for the synthetic charts (synthetic: Mann-Whitney $U = 131594$, $p < 0.001$; real-world: Mann-Whitney $U = 48599.5$, $p = 0.132$).

The results together suggest that we can accept H2 that states that including external information in the caption helps reinforce the message in the caption and users are more likely to take away from the feature described in the caption.

H2 Additional Results. Figure 7 (bottom row) breaks down the ratio of the takeaways that mention the feature described in the caption by level of prominence of the feature. The figure shows that there is usually an increase in the ratio of the takeaways that mentioned the feature described in the caption when the caption included external information for each level of prominence. Among the differences, we only found significant difference when the caption mentioned a non-prominent feature for synthetic charts ($Z = 3.027$, $p = 0.001$). Further study could shed light on the correlation between the prominence of the feature described in the caption and how external information affects the readers’ takeaways.

5 DESIGN GUIDELINES

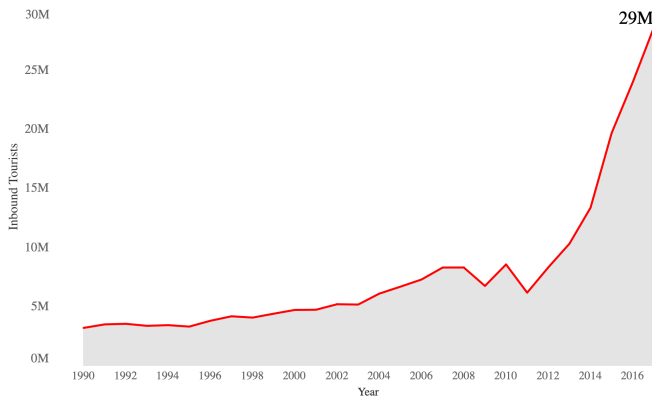
Our findings indicate that the readers will take away from the feature doubly emphasized by both the chart and caption if they provide a coherent message. However, when the chart and caption diverge in terms of the feature that they are emphasizing, readers are less likely to use information from the caption in their takeaways. To improve the efficacy of the chart-caption pair, authors could (1) design the chart to make the feature described in the caption more prominent and (2) include external information in the caption to give more context to the information in the caption.

There are several ways for authors to emphasize aspects of the data in a chart so that readers’ attention is drawn to these visual features. One technique is to ensure that aspects of the data such as trends and outliers are presented at the right level of detail or interval range; too-broad of a measurement interval may hide a signal. For example, assume that we were given the chart in Figure 8a with the caption in Figure 8b. The decrease in 2009 is not very prominent because the large increase starting in 2011 overshadows the decrease. Zooming closer to the intended feature and cropping out irrelevant features (Figure 8b), helps make the feature more visually prominent. However, when zooming into the data in this manner, authors must take precaution to avoid removing important information or rendering the chart misleading [33, 35].

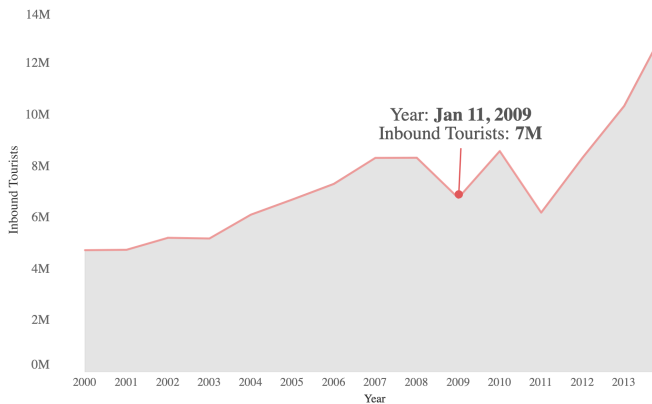
A simple way to further facilitate effective chart reading is to enhance the visualization with highlighting and overlays such as annotations to guide the audience’s attention to the image area they are describing [22] (Figure 8b). Sometimes, a different chart altogether may be more effective to emphasize a particular aspect of the data. For example, converting continuous data in line charts into discrete values could help emphasize individual values that the author would like to focus on. The consistency between the redesigned chart-caption pairs helps readers take away from the doubly emphasized feature (Figure 8).

6 FUTURE WORK

Chart and caption authoring tools. We would like to explore how this work can provide interesting implications for *both* chart



(a) “The cheap Yen and PM Abe’s tourism policy caused the number of tourists in Japan to steeply rise between 2011 and 2018.”



(b) “Due to the 2008 Financial Crisis, the number of tourists in Japan decreased in 2009.”

Figure 8: Examples of chart-caption pairs authored to emphasize the same feature in the data. (a) Both the caption and chart emphasize the sharp positive trend. (b) The original chart is modified to zoom into a portion of the time range and the feature is made more visually prominent with an annotation showing the dip in the number of tourists. The caption describes that dip with additional context.

and caption design to help the author effectively convey a specific point of view. Enhancements to visualization authoring tools could suggest chart design alternatives given a feature that the author would like to emphasize. Specifically, the system could go further by emphasizing features in the chart according to the main message the author wants to convey by automatically adding annotations to the chart, adding highlights, and adjusting levels of detail so that the chart and the caption deliver a concerted message. This will require formulating a high-level language specification that the authors can use to communicate to the system about their intents or a natural language processing module that can infer the authors’ intents based on the captions they write. Coordinating interaction between the chart and the caption such that hovering over the text in the caption would highlight the corresponding visual feature in the chart and vice-versa, is another interesting direction to pursue to help the reader. The resulting system would be a significant extension of the *interactive document reader* presented by Kong et

al. and Kim et al. [19, 23]. On the captioning side, a system could classify basic captions, captions about high-prominence features, and captions about low-prominence features. Based on the classification, the system could suggest external information to further emphasize the information presented.

Further exploration of caption variations. In this work, we use a template-based approach for generating captions to minimize the effect of the variation of natural language and to keep the experiment size reasonable. Simultaneously, we carefully vary the visual feature described in the caption and the presence of external information to best understand how people read captions and charts together to form their takeaways. Future work could study captions with various natural language expressions and different ways of emphasis. It would be useful to understand whether the relationship between multiple features in a caption (e.g., a simple list - “*There were major dips in employment in 2008 and 2020.*” or a comparison - “*The dip in 2020 was greater than the dip in 2008.*”) has an effect on what readers take away. Studying how our findings generalize to other types of external information (e.g., extrapolation, breakdown into subcategories) would be an interesting direction to pursue.

Generalization to other chart types. Our work explores how readers take away information when presented with univariate line charts and captions. Basic chart types still have prominent features (e.g., extrema in bar charts, outliers in scatterplots) and less prominent features (e.g., a point in a cluster in scatterplots). We expect similar findings would hold for those other chart types. We leave it to future work to confirm this intuition.

7 CONCLUSION

In this paper, we examine what readers take away from both a chart and its caption. Our results suggest that when the caption mentions visual features of differing prominence levels, the takeaways differ. When the caption mentions a specific feature, the takeaways also tend to mention that feature. We also observed that when a caption mentions a visually prominent feature, the takeaways more consistently mention that feature. On the other hand, when the caption mentions a less prominent feature, the readers’ takeaways are more likely to mention the most prominent prominence features than the feature described in the caption. We also find that including external information in the caption makes the readers more likely to form their takeaways based on the feature described in the caption. From the results of our study, we propose guidelines to better design charts and captions together; using visual cues and alternative chart representations, visual features can be made more prominent and be further emphasized by their descriptions in the caption. Design implications from this work provide opportunities for the authoring of chart and caption pairs in visual analysis tools to effectively convey a specific point of view to the reader.

ACKNOWLEDGMENTS

The authors thank the Stanford University HCI Group and Tableau Research for their feedback during the development of the studies. The authors also thank the reviewers for their feedback during the review cycle. This work is supported by NSF award III-1714647.

REFERENCES

- [1] Amazon Mechanical Turk 2020. Amazon Mechanical Turk. <https://www.mturk.com>.
- [2] Michelle A Borkin, Azalea A Vo, Zoya Bylinskii, Phillip Isola, Shashank Sunkavalli, Aude Oliva, and Hanspeter Pfister. 2013. What makes a visualization memorable? *IEEE Transactions on Visualization and Computer Graphics* 19, 12 (2013), 2306–2315.
- [3] J. Bransford. 1979. *Human Cognition: Learning, Understanding, and Remembering*. Wadsworth Publishing Company. <https://books.google.com/books?id=7TKLnQEACAAJ>
- [4] Sandra Carberry, Stephanie Elzer, and Seniz Demir. 2006. Information Graphics: An Untapped Resource for Digital Libraries. In *Proceedings of the 29th Annual International ACM SIGIR Conference on Research and Development in Information Retrieval* (Seattle, Washington, USA) (SIGIR '06). Association for Computing Machinery, New York, NY, USA, 581–588. <https://doi.org/10.1145/1148170.1148270>
- [5] Stuart K. Card, Jock D. Mackinlay, and Ben Shneiderman (Eds.). 1999. *Readings in Information Visualization: Using Vision to Think*. Morgan Kaufmann Publishers Inc., San Francisco, CA, USA.
- [6] Charles Chen, Ruiyi Zhang, Sungchul Kim, Scott Cohen, Tong Yu, Ryan Rossi, and Razvan Bunescu. 2019. Neural caption generation over figures. In *Adjunct Proceedings of the 2019 ACM International Joint Conference on Pervasive and Ubiquitous Computing and Proceedings of the 2019 ACM International Symposium on Wearable Computers*. 482–485.
- [7] Charles Chen, Ruiyi Zhang, Eunye Koh, Sungchul Kim, Scott Cohen, Tong Yu, Ryan Rossi, and Razvan Bunescu. 2019. Figure captioning with reasoning and sequence-level training. *arXiv preprint arXiv:1906.02850* (2019).
- [8] Zhe Cui, Sriram Karthik Badam, Adil Yalçın, and Niklas Elmqvist. 2018. DataSite: Proactive Visual Data Exploration with Computation of Insight-based Recommendations. *CoRR abs/1802.08621* (2018). [arXiv:1802.08621](https://arxiv.org/abs/1802.08621) <http://arxiv.org/abs/1802.08621>
- [9] Çağatay Demiralp, Peter J Haas, Srinivasan Parthasarathy, and Tejaswini Pedapati. 2017. Foresight: Rapid data exploration through guideposts. *arXiv preprint arXiv:1709.10513* (2017).
- [10] H. Egeth and S. Yantis. 1997. Visual attention: control, representation, and time course. *Annual review of psychology* 48 (1997), 269–97.
- [11] Stephanie Elzer, Sandra Carberry, Daniel Chester, Seniz Demir, Nancy Green, Ingrid Zukerman, and Keith Trnka. 2005. Exploring and exploiting the limited utility of captions in recognizing intention in information graphics. In *ACL*. 223–230.
- [12] Stephanie Elzer, Sandra Carberry, and Ingrid Zukerman. 2011. The automated understanding of simple bar charts. *Artificial Intelligence* 175, 2 (2011), 526–555.
- [13] Massimo Fasciano and Guy Lapalme. 1996. Postgraphe: a system for the generation of statistical graphics and text. In *Eighth International Natural Language Generation Workshop*.
- [14] Leo Ferres, Petro Verkhoglyad, Gitte Lindgaard, Louis Boucher, Antoine Chretien, and Martin Lachance. 2007. Improving Accessibility to Statistical Graphs: The IGraph-Lite System. In *Proceedings of the 9th International ACM SIGACCESS Conference on Computers and Accessibility* (Tempe, Arizona, USA) (Assets '07). Association for Computing Machinery, New York, NY, USA, 67–74. <https://doi.org/10.1145/1296843.1296857>
- [15] Mary Hegarty and Marcel-Adam Just. 1993. Constructing mental models of machines from text and diagrams. *Journal of memory and language* 32, 6 (1993), 717–742.
- [16] Kevin Hu, Diana Orghian, and César Hidalgo. 2018. DIVE: A Mixed-Initiative System Supporting Integrated Data Exploration Workflows. In *Proceedings of the Workshop on Human-In-the-Loop Data Analytics*. ACM, 5.
- [17] Jessica Hullman, Nicholas Diakopoulos, and Eytan Adar. 2013. Contextifier: automatic generation of annotated stock visualizations. In *Proceedings of the SIGCHI Conference on Human Factors in Computing Systems*. 2707–2716.
- [18] Dae Hyun Kim, Enamul Hoque, and Maneesh Agrawala. 2020. Answering Questions about Charts and Generating Visual Explanations. In *Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems*. 1–13.
- [19] Dae Hyun Kim, Enamul Hoque, Juho Kim, and Maneesh Agrawala. 2018. Facilitating Document Reading by Linking Text and Tables. In *UIST (UIST '18)*. ACM, 423–434.
- [20] Ha-Kyung Kong, Zhicheng Liu, and Karrie Karahalios. 2018. Frames and slants in titles of visualizations on controversial topics. In *Proceedings of the 2018 CHI Conference on Human Factors in Computing Systems*. 1–12.
- [21] Ha-Kyung Kong, Zhicheng Liu, and Karrie Karahalios. 2019. Trust and Recall of Information across Varying Degrees of Title-Visualization Misalignment. In *Proceedings of the 2019 CHI Conference on Human Factors in Computing Systems* (Glasgow, Scotland UK) (CHI '19). Association for Computing Machinery, New York, NY, USA, 1–13. <https://doi.org/10.1145/3290605.3300576>
- [22] N. Kong and M. Agrawala. 2012. Graphical Overlays: Using Layered Elements to Aid Chart Reading. *IEEE Transactions on Visualization and Computer Graphics* 18, 12 (2012), 2631–2638.
- [23] Nicholas Kong, Marti A. Hearst, and Maneesh Agrawala. 2014. Extracting References Between Text and Charts via Crowdsourcing. In *Proceedings of the SIGCHI Conference on Human Factors in Computing Systems* (Toronto, Ontario, Canada) (CHI '14). ACM, New York, NY, USA, 31–40. <https://doi.org/10.1145/2556288.2557241>
- [24] Andrew Large, Jamshid Beheshti, Alain Breuleux, and Andre Renaud. 1995. Multimedia and comprehension: The relationship among text, animation, and captions. *Journal of the American society for information science* 46, 5 (1995), 340–347.
- [25] Pengyuan Li, Xiangying Jiang, and Hagit Shatkay. 2018. Extracting Figures and Captions from Scientific Publications. In *Proceedings of the 27th ACM International Conference on Information and Knowledge Management* (Torino, Italy) (CIKM '18). Association for Computing Machinery, New York, NY, USA, 1595–1598. <https://doi.org/10.1145/3269206.3269265>
- [26] J. Liang and M. L. Huang. 2010. Highlighting in Information Visualization: A Survey. In *2010 14th International Conference on Information Visualisation*. 79–85.
- [27] Shixia Liu, Michelle X. Zhou, Shimei Pan, Weihong Qian, Weijia Cai, and Xiaoxiao Lian. 2009. Interactive, Topic-Based Visual Text Summarization and Analysis. In *Proceedings of the 18th ACM Conference on Information and Knowledge Management* (Hong Kong, China) (CIKM '09). Association for Computing Machinery, New York, NY, USA, 543–552. <https://doi.org/10.1145/1645953.1646023>
- [28] Shixia Liu, Michelle X. Zhou, Shimei Pan, Yangqiu Song, Weihong Qian, Weijia Cai, and Xiaoxiao Lian. 2012. TIARA: Interactive, Topic-Based Visual Text Summarization and Analysis. *ACM Trans. Intell. Syst. Technol.* 3, 2, Article 25 (Feb. 2012), 28 pages. <https://doi.org/10.1145/2089094.2089101>
- [29] Laura E Matzen, Michael J Haas, Kristin M Divis, Zhiyuan Wang, and Andrew T Wilson. 2017. Data visualization saliency model: A tool for evaluating abstract data visualizations. *IEEE transactions on visualization and computer graphics* 24, 1 (2017), 563–573.
- [30] Microsoft Q & A 2020. Microsoft Q & A. <https://powerbi.microsoft.com/en-us/documentation/powerbi-service-q-and-a/>.
- [31] Vibhu Mittal, Steven Roth, Johanna Moore, Joe Mattis, and Giuseppe Carenini. 1995. Generating Explanatory Captions for Information Graphics. *Proceedings of the International Joint Conference on Artificial Intelligence*, 1276–1283.
- [32] Gwen C Nugent. 1983. Deaf students' learning from captioned instruction: The relationship between the visual and caption display. *The Journal of Special Education* 17, 2 (1983), 227–234.
- [33] Shaun O'Brien and Claire Lauer. 2018. Testing the susceptibility of users to deceptive data visualizations when paired with explanatory text. In *Proceedings of the 36th ACM International Conference on the Design of Communication*. 1–8.
- [34] A. Ottley, E. M. Peck, L. T. Harrison, D. Afergan, C. Ziemkiewicz, H. A. Taylor, P. K. J. Han, and R. Chang. 2016. Improving Bayesian Reasoning: The Effects of Phrasing, Visualization, and Spatial Ability. *IEEE Transactions on Visualization and Computer Graphics* 22, 1 (2016), 529–538.
- [35] Anshul Vikram Pandey, Katharina Rall, Margaret L Satterthwaite, Oded Nov, and Enrico Bertini. 2015. How deceptive are deceptive visualizations? An empirical analysis of common distortion techniques. In *Proceedings of the 33rd Annual ACM Conference on Human Factors in Computing Systems*. 1469–1478.
- [36] Pew Research 2020. Pew Research. <https://www.pewresearch.org>.
- [37] Xin Qian, Eunye Koh, Fan Du, Sungchul Kim, and Joel Chan. 2020. A Formative Study on Designing Accurate and Natural Figure Captioning Systems. In *Extended Abstracts of the 2020 CHI Conference*. 1–8.
- [38] A. Srinivasan, S. M. Drucker, A. Endert, and J. Stasko. 2019. Augmenting Visualizations with Interactive Data Facts to Facilitate Interpretation and Communication. *IEEE Transactions on Visualization and Computer Graphics* 25, 1 (2019), 672–681. <https://doi.org/10.1109/TVCG.2018.2865145>
- [39] Tableau Public 2020. Tableau Public. <https://public.tableau.com>.
- [40] Tableau Software 2020. Tableau Software. <http://www.tableau.com>.
- [41] Edward Tufte. 1990. *Envisioning Information*. Graphics Press, USA.
- [42] Manasi Vartak, Samuel Madden, and Aditya N Parmeswaran. 2015. SEEDB : Supporting Visual Analytics with Data-Driven Recommendations. (2015).
- [43] Washington Post 2020. Washington Post. <https://www.washingtonpost.com>.
- [44] Furu Wei, Shixia Liu, Yangqiu Song, Shimei Pan, Michelle X. Zhou, Weihong Qian, Lei Shi, Li Tan, and Qiang Zhang. 2010. TIARA: A Visual Exploratory Text Analytic System. In *Proceedings of the 16th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining* (Washington, DC, USA) (KDD '10). Association for Computing Machinery, New York, NY, USA, 153–162. <https://doi.org/10.1145/1835804.1835827>
- [45] Wikipedia 2020. Wikipedia. <https://www.wikipedia.org>.
- [46] Graham Wills and Leland Wilkinson. 2010. Autovis: automatic visualization. *Information Visualization* 9, 1 (2010), 47–69.
- [47] Cindy Xiong, Lisanne van Weelden, and Steven Franconeri. 2019. The curse of knowledge in visual data communication. *IEEE Transactions on Visualization and Computer Graphics* (2019).
- [48] Jin Yu, Ehud Reiter, Jim Hunter, and Somayajulu Sripada. 2003. SumTime-Turbine: A Knowledge-Based System to Communicate Gas Turbine Time-Series Data. In *Developments in Applied Artificial Intelligence*, Paul W. H. Chung, Chris Hinde, and Moonis Ali (Eds.). Springer Berlin Heidelberg, Berlin, Heidelberg, 379–384.