

Context-Aware Explanations for Spatialized Document Layouts

Wei Liu

Department of Computer Science, Virginia Tech
Blacksburg, Virginia, USA
wliu3@vt.edu

Chris North

Department of Computer Science, Virginia Tech
Blacksburg, Virginia, USA
north@vt.edu

John Wenskovitch

Department of Computer Science, Virginia Tech
Blacksburg, Virginia, USA
jw87@vt.edu

Rebecca Faust

Department of Computer Science, Tulane University
New Orleans, Louisiana, USA
rf Faust1@tulane.edu

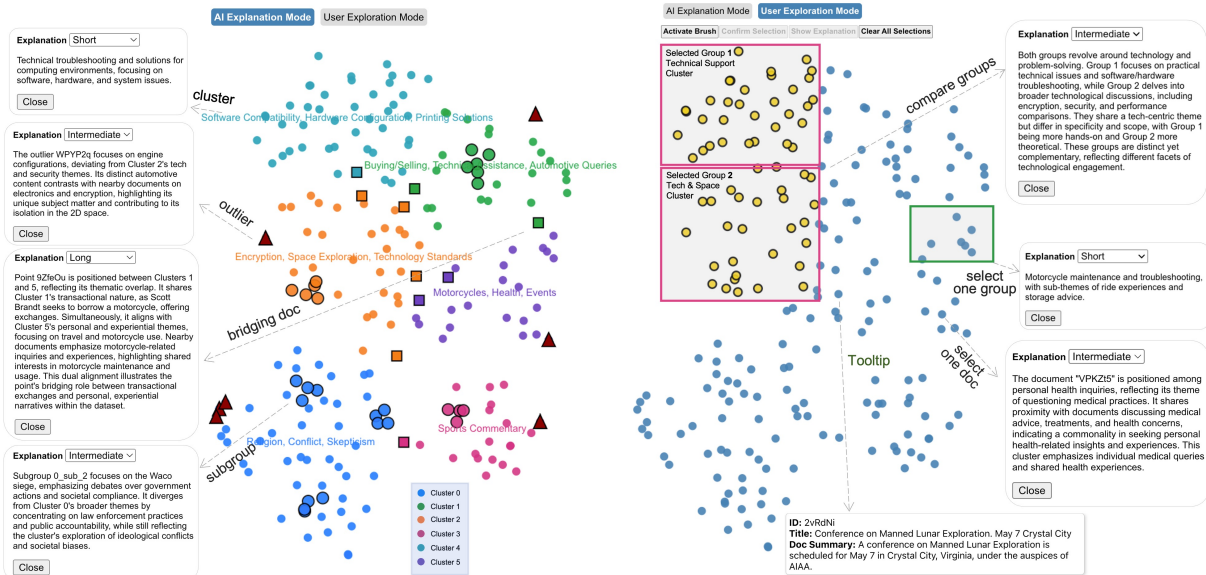


Figure 1: Context-aware explanations for spatialized document layouts, illustrated on a news dataset. In AI Explanation Mode (left), CAPE generates explanations at multiple levels of detail for salient spatial patterns, including clusters, subgroups, outliers, and bridging documents. In User Exploration Mode (right), users can inspect individual documents, request explanations for selected regions, and compare regions. Explanations integrate document semantics with spatial relationships to describe how regions and documents are organized within the layout.

Abstract

Spatialized document layouts are widely used for exploratory analysis of text corpora, but interpreting the spatial organization of documents and the relationships between regions remains challenging. Existing approaches primarily summarize document content or explain how layouts are generated, providing limited support for understanding spatial relationships within the layout itself. We present CAPE, a context-aware explanation framework that generates natural-language explanations grounded in both document semantics and layout-derived spatial context. CAPE identifies salient spatial patterns (e.g., clusters, subgroups, outliers, and bridging documents) and constructs multi-level contextual representations to guide LLM-based explanation generation. It supports both AI-guided overview and user-driven exploration, with explanations

available at multiple levels of detail. We demonstrate CAPE on news and scholarly document layouts and evaluate it in a controlled user study against keyword-based and content-only LLM baselines. Our results suggest that spatially grounded explanations are perceived as more helpful than content-only baselines for interpreting the spatial organization of document layouts.

CCS Concepts

• **Human-centered computing** → **Visual analytics**; *Information visualization*; *Visualization systems and tools*; • **Computing methodologies** → *Natural language generation*.

Keywords

Spatialized document layouts, Context-aware explanations, Document visualization, Visual analytics, Large language models

ACM Reference Format:

Wei Liu, John Wenskovich, Chris North, and Rebecca Faust. 2026. Context-Aware Explanations for Spatialized Document Layouts. In *Graphics Interface 2026 (GI '26)*, June 9–12, 2026, Kitchener-Waterloo, ON, Canada. ACM, New York, NY, USA, 10 pages. <https://doi.org/10.1145/nnnnnnn.nnnnnnn>

1 Introduction

Spatialized layouts of document collections are widely used to support exploratory analysis of text corpora, including workflows such as literature mapping, news exploration, and open-ended corpus sensemaking [9, 28]. By arranging documents in a two-dimensional (2D) space, these layouts allow users to inspect thematic groupings, compare regions, and explore relationships through spatial organization [11, 32, 46]. Such representations appear in many forms, including embedding-based projections, manually organized maps, and interactively refined layouts [4, 11, 19, 32]. Across these settings, users often treat the layout itself as an **interpretive workspace** for making sense of document collections [4, 8, 15].

Despite their usefulness, spatialized document layouts remain difficult to interpret [28, 35]. While the layout makes visible structures such as clusters, outliers, and boundary cases, it does not by itself explain how these structures should be understood in relation to document content. Users may see that documents form groups, that some regions lie close together, or that a document lies between multiple groups, but understanding the meaning of these spatial relationships requires inspecting many documents and mentally integrating local and global context. This process is effortful and time-consuming, especially in large heterogeneous collections.

Existing approaches provide only partial support for this interpretive challenge. Content-oriented techniques such as keyword summaries and topic labels describe what documents or regions are about, but do not explain how they relate within the spatial layout [23, 48]. Explainable AI methods, in contrast, often focus on how embeddings or projections are generated rather than on how users interpret the resulting spatial representation [5, 28]. More recently, large language models (LLMs) have enabled fluent summaries of document content, but these explanations are typically grounded in text alone and do not account for spatial relationships such as proximity, separation, or intermediate positioning [51]. As a result, current methods offer limited support for *interpreting the spatial organization of document collections*.

In this work, we introduce **CAPE (Context-Aware Explanations)**, a framework for generating natural-language explanations for spatialized document layouts. CAPE treats the current layout as the object of explanation and supports interpretation of its spatial organization, rather than explaining how the layout was produced. To do so, CAPE integrates semantic content with layout-derived spatial context at multiple levels, including global structure, local neighborhoods, and pattern-specific characteristics. This enables CAPE to generate explanations for salient spatial patterns such as clusters, subgroups, outliers, and bridging documents, as well as for user-selected documents and regions. The goal is to help users understand not only what documents or regions are about, but also how they relate to surrounding parts of the layout.

We target analysts and researchers who use spatialized document layouts as exploratory workspaces for understanding document

collections, such as news corpora, scholarly literature, and related text collections. This work is guided by two research questions:

(RQ1) How can natural-language explanations for spatialized document layouts be constructed by integrating semantic information with spatial context?

(RQ2) Does grounding explanations in spatial context help users interpret spatial relationships and organization within a layout, compared with content-only baselines?

To address these questions, we present the CAPE framework, demonstrate it through interactive visualizations, and evaluate it in a controlled user study against keyword-based and content-only LLM baselines. Our results indicate that spatially grounded explanations are perceived as more helpful for interpreting spatial organization and relationships within document layouts than content-only baselines.

This work makes the following contributions:

- A context-aware explanation framework for spatialized document layouts that integrates semantic and spatial context.
- A multi-level explanation strategy supporting both AI-guided overview and user-driven exploration.
- An interactive design demonstrating how context-aware explanations can be integrated into exploratory analysis of spatialized document layouts.
- Empirical evidence that spatially grounded explanations are perceived as more helpful than content-only baselines for interpreting spatialized document layouts.

2 Related Work

CAPE draws on prior work in spatialized document layouts, explanations of spatial representations, and language-based support for data interpretation.

Spatialized Document Layouts Spatialized document layouts arrange document collections in a 2D space, where spatial proximity reflects semantic relationships [4, 11]. Prior work has proposed a range of approaches for constructing and interacting with such layouts, including topic model visualizations, embedding-based projections, and interactive document maps [1, 4, 8, 11, 32, 50]. These methods allow users to identify clusters, compare regions, and navigate document collections through spatial organization [28]. To support interpretation, many methods augment layouts with content-oriented summaries such as keyword summaries, topic labels, or representative documents [9, 18, 23, 33, 40, 42, 48]. For example, systems such as TopicPanorama [48] and Scattertext [23] provide keyword- or phrase-based summaries to help users understand document collections. While these techniques help characterize what documents or regions are about, they primarily focus on semantic content, leaving visible spatial relationships within the layout, such as grouping, separation, and intermediate positioning, largely unexplained.

Explaining Spatial Representations Explainability has been extensively studied in the context of embeddings, dimensionality reduction, and visual representations [13, 30, 39, 49]. Prior work has proposed methods to explain projection behavior, feature contributions, and distortions in low-dimensional embeddings, helping users understand how high-dimensional data are mapped into spatial representations [5, 28]. These approaches are valuable for assessing

projection quality and diagnosing model behavior. However, their primary focus is on explaining the computational processes that generate a layout, leaving users to interpret the resulting spatial organization largely on their own. In many exploratory settings, users reason directly about spatial relationships such as proximity, separation, grouping, and overlap, without reference to the underlying model. As a result, methods that explain projection mechanisms do not necessarily support understanding of how documents and regions should be interpreted within the layout. CAPE complements this line of work by focusing on the layout as an interpretive representation and generating explanations grounded in layout-derived spatial context.

Language-Based Explanations for Data Interpretation Natural language explanations are increasingly used to support data interpretation, including summarizing document collections, labeling clusters, and assisting exploratory analysis workflows [7, 45, 51]. Recent advances in LLMs have enabled fluent and flexible explanations based on textual content, and have been integrated into visualization workflows to support sensemaking [3, 14, 44]. For example, methods such as LangLasso [6] and TextCluster Explainer [38] allow users to obtain textual summaries of selected regions or clusters, improving accessibility and supporting interactive exploration. However, these approaches primarily describe the semantic content of documents or regions and do not explicitly account for their relationships within the spatial organization of the layout.

Summary and Gap Across these lines of work, existing approaches either focus on summarizing document content or explaining how spatial representations are generated. *Few integrate semantic content with layout-derived spatial context to support understanding of how documents, regions, and spatial patterns relate within the layout.* This gap motivates the design of CAPE.

3 Design Considerations for Explaining Spatialized Document Layouts

Spatialized document layouts serve as exploratory workspaces [1, 9] in which users interpret document collections by reasoning about visible structures such as proximity, separation, grouping, and boundary cases [1, 10, 43]. Prior work on sensemaking and visual analytics suggests that users construct meaning by iteratively comparing regions, identifying patterns, and relating local observations to global structure [1, 2, 10, 12]. Interpreting a spatialized layout is therefore a relational and iterative task centered on understanding how documents, regions, and patterns are organized and connected within the layout. This perspective motivates four design considerations for explaining spatial organization in document layouts.

DC1: Integrating Spatial and Semantic Context. Interpretation requires connecting what documents are about with how they are organized in the layout. Spatial proximity is often interpreted as similarity, separation suggests divergence, and intermediate positioning may indicate overlap between themes [1, 2, 10, 12, 43]. Explanations that consider only semantic content or only spatial structure provide incomplete support for interpretation.

DC2: Explaining Salient Spatial Patterns. Users tend to organize their understanding around perceptually salient structures such as clusters, subgroups, outliers, and bridging cases [1, 2, 10, 12,

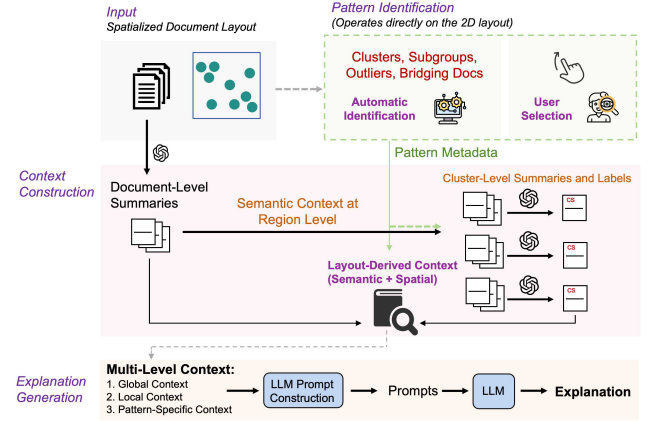


Figure 2: Overview of the CAPE framework. CAPE identifies explanation targets through automatic detection of salient spatial patterns (clusters, subgroups, outliers, and bridging documents) or user selection of documents or regions. It then integrates document semantics with layout-derived spatial context to guide LLM-based explanation generation.

26, 43]. These patterns serve as natural units for reasoning about the layout [8, 16, 21, 36]. Explanations should therefore target not only individual documents but also higher-level structures and their relationships.

DC3: Supporting Multiple Levels of Explanation. Exploratory analysis involves moving between high-level overview and detailed inspection [9, 17, 22, 44]. Users often begin with concise summaries to understand global structure, and then require more detailed explanations when examining specific regions or documents. Explanation mechanisms should therefore support multiple levels of detail.

DC4: Enabling User-Driven Exploration. Sensemaking is iterative and user-driven, with users shifting focus based on emerging questions [22, 37, 44]. Users may wish to request explanations for specific documents, compare regions, or examine boundary cases. Explanations should therefore be generated on demand and adapt to user-selected areas of interest.

4 CAPE Framework

CAPE is a framework for generating context-aware explanations of spatialized document layouts. Given a 2D layout of documents and their associated textual content, CAPE produces natural-language explanations that integrate semantic meaning with layout-derived spatial relationships.

To address **RQ1** and operationalize the design considerations described in Section 3, CAPE follows a structured pipeline consisting of three stages (Figure 2): (1) **identifying salient spatial patterns** in the layout (DC2, DC4), (2) **constructing structured contextual representations** that integrate semantic and spatial information (DC1), and (3) **generating explanations** conditioned on the pattern type or user query (DC2, DC4), with outputs available at multiple levels of detail (DC3).

4.1 Input: Spatialized Document Layouts

CAPE operates on a spatialized document layout in which documents are arranged in a 2D space and associated with textual representations. **The framework is agnostic to how the layout is generated**, and can be applied to embedding-based projections, manually organized maps, or interactively refined layouts.

CAPE treats the layout as the current representation presented to the user and generates explanations with respect to this representation. CAPE does not assume that spatial proximity strictly reflects high-dimensional similarity. Instead, it focuses on spatial relationships as they appear in the layout, including proximity, separation, and intermediate positioning. For dynamic or interactive layouts, CAPE can be applied to the current state of the layout, with contextual representations recomputed as spatial relationships change.

4.2 Identifying Salient Spatial Patterns

To support interpretation, CAPE identifies spatial patterns that are likely to draw users' attention during exploration (*DC2*). These patterns correspond to commonly observed and perceptually salient structures in spatialized document layouts, including:

- **Clusters**: cohesive regions of related documents;
- **Subgroups**: denser local structures within clusters;
- **Outliers**: documents that are spatially separated from surrounding points;
- **Bridging documents**: documents positioned between multiple regions and reflecting potential thematic overlap.

This stage is designed to approximate how users visually segment and interpret spatial layouts, not to provide definitive or optimal pattern detection. The pattern identification component is therefore modular and can incorporate different spatial analysis techniques. In our implementation, clusters are approximated using broad spatial groupings of documents, while subgroups are identified based on local density within clusters. Outliers are detected based on spatial separation from surrounding regions, and bridging documents are identified through proximity to multiple clusters.

Each identified pattern is associated with basic spatial metadata, including its location in the layout, its relationship to neighboring regions, and the documents it contains and those nearby. This information is not presented directly as algorithmic output, but serves as input for constructing contextual representations used in explanation generation.

In addition to automatically identified patterns, CAPE supports user-driven selection of documents or regions (*DC4*). Explanations can therefore be generated both for automatically identified structures and for user-defined areas of interest. The framework is selection-agnostic, allowing different interaction techniques such as rectangular or lasso-based selections.

4.3 Structured Context Construction and Explanation Generation

To generate explanations that reflect both document semantics and spatial organization (*DC1*), CAPE constructs a **structured contextual representation** for each explanation target and uses this

representation to guide explanation generation. Details of the contextual representation and the prompt design used for explanation generation are provided in the supplementary materials.

Providing full document collections to a language model is inefficient and can introduce noise. CAPE therefore constructs a compact representation that captures the contextual information most relevant to each spatial pattern or user-defined selection. This enables explanations to remain both informative and scalable, while grounding them in the structure of the layout.

4.3.1 Context Representation. CAPE represents each explanation target using a combination of semantic abstractions and spatial descriptors.

At the **semantic level**, each document is represented using a concise summary that captures its main idea. These summaries are precomputed and reused across explanation requests. For higher-level structures such as clusters, CAPE derives aggregate representations (e.g., representative summaries) that characterize dominant themes within a region.

At the **spatial level**, CAPE encodes relationships derived from the layout itself. For each pattern or selected region, this includes:

- its position within the layout,
- nearby documents in the 2D space, and
- relationships to surrounding regions or clusters.

Additional spatial descriptors are incorporated depending on the pattern type. For example, outliers are characterized by relative isolation from surrounding regions, while bridging documents are associated with proximity to multiple regions. These semantic and spatial components form a compact contextual representation that captures **both what documents are about and how they are organized relative to one another in the layout**.

4.3.2 Multi-Level Context Integration. To support interpretation at different levels of granularity, CAPE organizes contextual information into three complementary levels:

- **Global context**, which situates the target pattern within the overall layout structure, including region-level summaries and relationships to neighboring regions;
- **Local context**, which captures nearby documents and their summaries, reflecting similarity, contrast, and transitions across adjacent areas;
- **Pattern-specific context**, which captures structural characteristics associated with different spatial patterns, such as relative isolation, proximity to multiple regions, or internal relationships within a cluster.

This multi-level organization allows CAPE to balance breadth and specificity. Global context provides structural overview, local context captures fine-grained relationships, and pattern-specific context focuses on characteristics that are particularly relevant for interpreting a given pattern or selection.

4.3.3 Pattern- and Query-Aware Explanation Generation. Given a structured contextual representation, CAPE generates explanations by conditioning a language model on both the contextual information and the explanatory intent. In our implementation, explanations were generated using the OpenAI API with GPT-4o, with the temperature set to 0.2 [34]. Explanation generation is pattern-

and query-aware: CAPE **tailors prompts to each explanation target**. For example:

- For **clusters**, explanations emphasize shared themes and distinctions from neighboring regions;
- For **subgroups**, explanations describe how local variations relate to the broader cluster;
- For **outliers**, explanations focus on how a document differs from and relates to surrounding regions;
- For **bridging documents**, explanations highlight connections across multiple regions.

CAPE also supports user-driven queries. When users select individual documents, explanations relate the document’s content to its surrounding context. When users select groups of documents or regions, CAPE produces summaries of shared properties or comparative explanations that highlight similarities and differences.

This pattern- and query-aware strategy enables explanations to align with users’ analytical intent while remaining grounded in the spatial structure of the layout.

4.3.4 Multi-Level Explanation Outputs. To accommodate different interpretive needs, CAPE produces explanations at three levels of detail (DC3): **short, intermediate, and long**. Short explanations provide concise summaries for rapid orientation. Intermediate explanations include additional contextual detail about nearby regions and spatial relationships. Long explanations further elaborate on semantic variation and broader structural context. All three levels are generated from the same underlying contextual representation, differing only in the amount of information expressed in the output. This design supports progressive disclosure of information, allowing users to balance readability and informational depth during exploration.

5 Interaction Design for Context-Aware Explanations

To operationalize the design considerations described in Section 3, we illustrate how CAPE supports context-aware explanations in spatialized document layouts through two complementary interaction modes (Figure 1): an AI-guided overview mode (Figure 1, left) for rapid orientation and a user-driven exploration mode (Figure 1, right) for on-demand analysis. These modes, together with multi-level explanations, enable users to interpret spatial organization at different stages of exploratory analysis. Across both modes, users can hover over documents to inspect metadata such as titles and summaries, and click on highlighted patterns or user-defined selections to view their associated explanations.

5.1 AI-Guided Overview Mode

The AI-guided overview mode provides an initial interpretive structure by automatically highlighting salient spatial patterns within the layout, including clusters, subgroups, outliers, and bridging documents (DC2). This design foregrounds perceptually salient structures during early-stage sensemaking, allowing users to quickly identify meaningful regions and relationships.

Each highlighted pattern is associated with a context-aware explanation that integrates semantic summaries with spatial relationships. These explanations describe both the main themes of the

highlighted structure and how it relates to surrounding areas, helping users form an overview of the layout, identify areas of interest, and locate boundary or transitional cases for further analysis.

5.2 User-Driven Exploration Mode

The user-driven exploration mode generates explanations in response to user-selected documents or regions, supporting flexible and hypothesis-driven analysis (DC4).

Users can select individual documents to request explanations of how they are positioned within the layout. These explanations relate the document’s content to its surrounding context, including nearby documents and adjacent regions. When users select groups of documents or regions, they can obtain summaries of shared properties or comparative explanations that highlight similarities and differences across selections. This interaction mode supports a range of exploratory tasks, such as interpreting boundary cases, comparing neighboring regions, and examining localized variations.

5.3 Multi-Level Explanations Across Modes

Both interaction modes provide explanation outputs at multiple levels of detail (DC3), enabling users to balance rapid orientation with in-depth analysis. This progressive disclosure lets users start with concise summaries and expand to richer explanations as needed.

6 Usage Scenarios

We present two usage scenarios to illustrate how CAPE supports interpretation of spatialized document layouts in practice.

6.1 Exploring Themes in News Documents

We first demonstrate CAPE on a spatialized layout of 216 documents sampled from the 20 Newsgroups dataset, covering diverse topics including technology, politics, religion, science, and recreational activities [24]. Documents were embedded using OpenAI’s text-embedding-3-small model [34] and projected into a 2D space using t-SNE [47]. This scenario illustrates how CAPE supports interpretation of (1) global thematic structure, (2) variation within regions, (3) relationships between neighboring regions, and (4) boundary and transitional cases (Figure 1).

Using the AI-guided overview mode (Figure 1 left), CAPE highlights prominent regions and provides cluster-level explanations that summarize dominant themes while situating each region relative to surrounding areas. This supports rapid understanding of the overall structure of the layout and how major thematic areas are organized.

Within a large region centered on religious and ideological discussions (Figure 1 left-bottom), CAPE reveals finer-grained structure through subgroup-level explanations. These explanations distinguish localized perspectives, such as theological contradictions, societal attitudes, and debates about governmental authority, while relating them to the broader region. This helps users recognize meaningful variation within an otherwise cohesive area.

CAPE enables comparison between nearby regions by generating explanations that highlight differences in thematic focus. For example, two adjacent regions both relate to technology (Figure 1 right-top), but differ in scope: one emphasizes practical technical support and troubleshooting, while the other focuses on broader

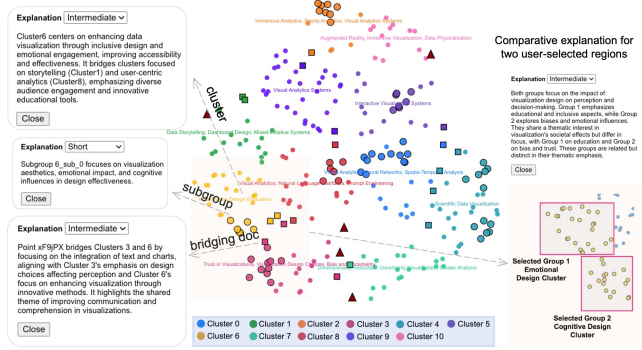


Figure 3: Context-aware explanations on a spatialized layout of IEEE VIS papers. CAPE explains salient spatial patterns identified in the 2D layout, with explanations available at multiple levels of detail (left). Users can also compare selected regions, as illustrated by the contrast between two user-selected design-focused groups (right).

topics such as encryption, security, and space exploration. Comparative explanations clarify how these regions are positioned close together while remaining distinct.

CAPE also supports interpretation of boundary cases by relating documents to multiple surrounding regions. For example, one bridging document (Figure 1 left) discusses borrowing a motorcycle for a trip, and is explained as combining transactional themes (associated with buying and selling activities) and more personal and experiential themes such as travel and motorcycle use, reflecting its intermediate position in the layout.

6.2 Exploring Trends in Scholarly Collections

We next apply CAPE to a spatialized layout of 252 papers from the IEEE VIS 2022 and 2023 proceedings [20]. Each paper is represented by its title and abstract, embedded using OpenAI’s text-embedding-3-small model [34], and projected into a 2D space using UMAP [31] (Figure 3). Compared to the news dataset, this collection reflects a more complex domain in which research themes, methods, and applications are often intertwined.

Using the AI-guided overview mode, CAPE identifies major research areas and provides cluster-level explanations that characterize dominant themes while situating each region within the broader layout. For example, CAPE describes a cluster (Figure 3 left) centered on inclusive and emotional design in terms of accessibility, engagement, and educational tools, and connects it to neighboring regions focused on storytelling and user-centric analytics.

CAPE further reveals variation within regions through subgroup-level explanations. Within the same cluster, subgroups (Figure 3 left-middle) highlight more specific thematic emphases, such as visualization aesthetics and cognitive influences on design effectiveness, helping distinguish localized structure within a broader research area.

CAPE also supports comparison between neighboring regions. As shown in Figure 3 (right), two user-selected groups are both related to visualization design and decision-making, but differ in focus: one centers on educational and inclusive aspects, while the

other emphasizes bias, trust, and emotional influences. These explanations clarify how regions can be closely positioned while maintaining distinct thematic orientations.

Finally, CAPE supports interpretation of boundary cases. For example, a bridging document (Figure 3 left-bottom) is explained as connecting design-oriented concerns about perception and cognition with methodological approaches for improving visualization techniques, highlighting its role in linking adjacent research areas.

7 User Study: Evaluating Perceived Quality of Explanations

To address RQ2, we conducted a controlled user study to evaluate the perceived quality of explanations generated under different conditions.

7.1 Study Design

We designed a **within-subjects** study to compare three explanation strategies:

(1) **CAPE**. Explanations generated using CAPE incorporated spatial context at multiple levels, including global layout structure, local neighborhoods, and pattern-specific information. All CAPE explanations were presented at an intermediate level of detail to ensure comparability.

(2) **LLM (Content Only)**. This baseline used the same language model as CAPE but was provided only with document content summaries, without spatial or neighborhood context. Explanation length was controlled to be comparable to CAPE outputs.

(3) **Keyword-based**. This baseline represents commonly used content-oriented labeling approaches in document visualization. For each spatial pattern, we extracted the top seven TF-IDF [41] keywords from associated documents.

Both CAPE and the LLM (Content Only) baseline used GPT-4o (temperature = 0.2) via the OpenAI API [34]. This design helps isolate the contribution of spatial context from differences in model capability. Table 1 provides a representative example of the explanations shown to participants under each condition, illustrating how the three strategies differed in the information they conveyed.

We evaluated explanations across four types of spatial patterns commonly observed in document layouts: clusters, subgroups, outliers, and bridging documents. Two instances of each pattern type were selected from the news layout described in Section 6.1, resulting in eight evaluation cases per participant.

7.2 Participants

We recruited 10 participants from a university mailing list. The sample included 6 PhD students, 3 master’s students, and 1 undergraduate student. Participants primarily came from computing-related backgrounds. Because background was collected as a multi-select question, categories overlap: 9 participants selected Computer Science, 2 selected Data Science, 2 selected Visualization, and 1 selected Human–Computer Interaction (HCI). Participants reported moderate to high familiarity with data visualization (7-point scale: range = 3–7, median = 5.5). All participants reported at least basic familiarity with interpreting document visualizations. This initial evaluation intentionally recruited participants with sufficient technical background to engage with spatialized document layouts and

Table 1: Representative example of explanations presented in the study.

Method	Explanation	Key Characteristics
CAPE (Intermediate length)	“This isolated document focuses on engine configurations, deviating from the broader Technology and Innovation Cluster. Its distinct automotive content contrasts with nearby documents on electronics and encryption, highlighting its unique topic and contributing to its isolation in the 2D space.”	Combines semantic and spatial context; explains how the document relates to surrounding regions within the layout.
LLM (Content Only) (Intermediate length)	“The document discusses the differences between flat-four (boxer) and inline-four (I4) engine configurations, highlighting the balance and mounting advantages of the flat-four used by Subaru and Porsche. It notes that flat-fours are smoother and more compact, making them suitable for low hoods and rear compartments, while I4s are better for transverse mounting.”	Fluent topical summary but does not account for spatial relationships or layout context.
Keyword-based	“flat, order, crank, second, engines, second order, subaru”	Lists frequent terms; no relational context.

explanation comparison tasks. The study was approved by the institutional review board.

7.3 Materials and Procedure

The study was conducted in person, with each session lasting approximately 25–35 minutes. Participants interacted with a setup that allowed them to view the spatialized document layout and corresponding explanations side by side. When inspecting the layout, participants could view document titles and summaries on demand. Two counterbalanced versions of the study materials were prepared. These versions rotated both (1) the mapping between explanation conditions and labels (A–C) and (2) the order of the eight evaluation patterns. After a brief introduction, participants completed a short tutorial and one practice trial. In each trial, participants were presented with a visualization and three corresponding explanations (one per condition). For each explanation, participants provided ratings on multiple quality dimensions (Section 7.4). After evaluating all three explanations for a given pattern, participants ranked them from most to least helpful. Each participant evaluated eight patterns. After the trials, participants completed a short questionnaire and provided open-ended feedback on helpful and unhelpful aspects of the explanations.

7.4 Measures

We collected participants’ subjective ratings using 7-point Likert scales (1 = Not at all, 7 = Very well) across five dimensions.

Clarity: how clear and easy the explanation is to understand.

Relevance: how well the explanation reflects the main themes of the documents or region.

Specificity: how detailed and informative the explanation is.

Position-awareness: how well the explanation helps participants understand how a document or region is positioned within the layout and how it relates to surrounding areas.

Usefulness: overall perceived helpfulness of the explanation for interpreting the spatial pattern.

We selected clarity, relevance, specificity, and usefulness because they are commonly used criteria in prior work evaluating explanations and visualization support tools, and introduced position-awareness as a study-specific dimension to capture support for understanding spatial relationships within the layout [25]. In addition to Likert-scale ratings and ranking data, we also collected

open-ended responses at the end of the study to capture qualitative feedback about explanation quality.

7.5 Data Analysis

All analyses were conducted at the participant level. For each participant, ratings were first averaged across the eight evaluated patterns to obtain a single score per condition and evaluation dimension. This aggregation treated the participant as the unit of analysis in the within-subjects design. Given the ordinal nature of Likert-scale data and the small sample size, we used non-parametric statistical tests. Descriptive statistics are reported as the median and interquartile range (IQR) of participant-level aggregated scores. To assess differences between conditions, we conducted Friedman tests for each evaluation dimension. When significant effects were observed, pairwise comparisons were performed using Wilcoxon signed-rank tests with Holm correction for multiple comparisons. Statistical significance was assessed at $\alpha = .05$. Effect sizes are reported using Kendall’s W for Friedman tests and rank-biserial correlation (r_{rb}) for Wilcoxon tests. For ranking data, participant-level average ranks were computed across the eight evaluated patterns and analyzed using the same procedure (Friedman tests followed by Holm-corrected Wilcoxon tests). Open-ended responses were analyzed using open coding. One researcher first assigned descriptive codes to all responses, after which a second researcher reviewed the coding and emerging themes. Discrepancies were discussed and resolved to consolidate a set of recurring themes.

7.6 Results

7.6.1 Rating Results. Explanation condition had a significant effect on all five evaluation dimensions (Friedman tests, all $p < .01$), with medium to large effect sizes: Clarity ($\chi^2(2) = 10.32, p = .0058, W = .52$), Relevance ($\chi^2(2) = 15.85, p < .001, W = .79$), Specificity ($\chi^2(2) = 15.85, p < .001, W = .79$), Position-awareness ($\chi^2(2) = 16.77, p < .001, W = .84$), and Usefulness ($\chi^2(2) = 17.59, p < .001, W = .88$). As shown in Figure 4, CAPE received consistently high ratings across all dimensions, whereas the keyword baseline received low ratings. The LLM (Content Only) condition received ratings similar to those of CAPE for Clarity, Relevance, and Specificity, but lower ratings for Position-awareness and Usefulness. Holm-corrected Wilcoxon signed-rank tests indicated that both CAPE and

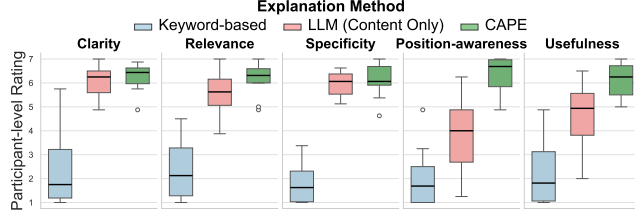


Figure 4: Participant ratings (7-point Likert scale) across explanation methods. CAPE receives higher ratings for position-awareness and usefulness, while ratings for clarity, relevance, and specificity are comparable to the LLM (Content Only) baseline. The keyword-based baseline receives the lowest ratings across all dimensions.

LLM (Content Only) were rated significantly higher than the keyword baseline across all dimensions (all $p_{\text{Holm}} < .05$). Comparing the two LLM-based conditions, CAPE received significantly higher ratings for Position-awareness ($W = 1.0$, $p_{\text{Holm}} = .0078$, $r_{rb} = .96$) and Usefulness ($W = 3.0$, $p_{\text{Holm}} = .0152$, $r_{rb} = .89$), with no significant differences observed for Clarity, Relevance, or Specificity. These results suggest that spatial grounding primarily improves perceived support for understanding spatial relationships, while both LLM-based conditions performed similarly on general explanation qualities.

7.6.2 Ranking Results. The ranking results were consistent with the rating results. Across all 80 trials (10 participants $\times$ 8 patterns), CAPE was ranked as the most helpful explanation in **64 out of 80 cases**, followed by LLM (Content Only) in 16 cases, and the Keyword-based baseline in 0 cases. Participant-level analysis showed a significant effect of explanation condition ($\chi^2(2) = 19.54$, $p < .001$, $W = .98$). Holm-corrected Wilcoxon signed-rank tests indicated that **CAPE was ranked significantly higher** than both LLM (Content Only) ($W = 0.0$, $p_{\text{Holm}} = .0071$, $r_{rb} = 1.00$) and the Keyword-based baseline ($W = 0.0$, $p_{\text{Holm}} = .0059$, $r_{rb} = 1.00$). LLM (Content Only) also outperformed the Keyword-based baseline ($W = 0.0$, $p_{\text{Holm}} = .0059$, $r_{rb} = 1.00$).

7.6.3 Qualitative Feedback. All participants provided open-ended comments, which offer additional insight into how different explanation strategies supported interpretation. Open coding produced four themes that organize the discussion below.

- *Combining semantic summaries with spatial reasoning:* connecting what a document or region is about with its spatial context in the layout.
- *Understanding relationships to nearby documents and regions:* describing how a document or pattern relates to neighbors.
- *Clarity and level of detail:* expressing the main idea clearly and at an appropriate level of detail.
- *Limitations of keyword-based explanations:* keyword-only explanations perceived as too vague to convey meaning or to clarify relationships within the layout.

Combining Semantic Summaries with Spatial Reasoning.

Many participants described helpful explanations as integrating both topical summaries and spatial context (7/10). For example, a

participant commented that helpful explanations “talk about the document and also why a document is placed where it’s placed” (P01), while another emphasized explanations that “summarized both the main topic and why it was positioned in a specific location” (P02).

Understanding Relationships to Nearby Documents and Regions. Several participants highlighted the importance of describing relationships to nearby documents or clusters (5/10). For example, a participant noted that helpful explanations should “[show] their relation to nearby documents” (P04). Another stated that “the document-cluster relationship[s] are what is difficult to find by hand” (P06). These responses suggest that spatially grounded explanations may make such relationships easier to interpret.

Clarity and Level of Detail. Participants also valued explanations that were clear and concise. Some participants described helpful explanations as easy to read and well-structured (4/10), such as “short but capturing the main theme with easy words” (P07) and “simple and only used jargon if necessary” (P10). One participant also noted that overly long explanations could be difficult to process, requiring them to “go back and forth and consume all information” (P09). These responses suggest a need to balance informativeness with readability.

Limitations of Keyword-Based Explanations. Participants frequently described keyword-based explanations as vague and difficult to interpret (7/10). For example, one participant noted that they “only gave several words... I have to guess the topic” (P02), while another described them as “completely unhelpful” because “the missing sentences made it hard to interpret the words properly” (P05). In addition, several participants mentioned that unhelpful explanations often failed to clarify relationships, making them less useful for understanding spatial organization (7/10).

7.6.4 Summary of Findings. Across ratings, rankings, and qualitative feedback, CAPE was consistently preferred over the keyword baseline and was most frequently ranked as the most helpful explanation. Compared to the content-only LLM baseline, CAPE was associated with higher perceived position-awareness and usefulness, while maintaining comparable clarity, relevance, and specificity. These findings suggest that incorporating spatial context supports users in interpreting how documents and regions relate within the layout.

8 Discussion

Explaining Spatial Layouts as Interpretive Representations. Spatialized document layouts are often treated as outputs of embedding and projection pipelines, with interpretability work focusing on how such layouts are generated [19, 21, 28]. In contrast, CAPE treats the layout itself as an *interpretive representation that users reason about during exploratory analysis*. In this setting, users often attend to visible spatial cues such as grouping, separation, and intermediate positioning, regardless of how the layout was produced. CAPE therefore complements prior work on explainable embeddings and dimensionality reduction by focusing on user-facing interpretation of visible spatial organization.

What Spatial Context Adds Beyond Content-Only Explanations. Our results suggest that the main contribution of spatial context lies not in improving linguistic quality, but in improving

support for reasoning about relationships within the layout. CAPE and the content-only LLM baseline received comparable ratings for clarity, relevance, and specificity, but CAPE was rated higher for position-awareness and overall usefulness. This pattern suggests that content-only explanations are effective for summarizing topical content, while spatially grounded explanations provide *additional support for understanding documents and regions in relation to surrounding spatial context*. This is consistent with prior work describing sensemaking in spatial layouts as a relational process centered on comparison, proximity, separation, and transition [1, 10, 43].

Implications for Explanation Design. These findings suggest three implications for explanation design in spatial analytic workflows. First, explanations should explicitly address relationships between documents and regions, not just summarize topical content. Second, explanations may benefit from adapting to different spatial pattern types, since clusters, subgroups, outliers, and bridging cases involve different interpretive questions. Third, explanation mechanisms should support both rapid overview and on-demand exploration, ideally with multiple levels of detail, so that users can move between broad orientation and focused inspection.

Scope and Future Directions. CAPE is designed to explain layouts as presented to users, not to validate projection faithfulness or correct distortions introduced during layout generation. This makes the framework applicable across embedding-based, manually organized, and interactively refined layouts, but also means that explanations inherit biases present in the current representation [11, 27, 32, 46]. Future work could explore adaptive explanation mechanisms, integration of additional metadata such as temporal or categorical information, and extension to other forms of spatialized analytic representations.

Limitations. This work has several limitations. First, the user study involved a relatively small sample ($N = 10$) drawn primarily from computing-related backgrounds, which limits generalizability beyond technically experienced users. We also did not collect detailed demographic information such as gender; prior work suggests that individual differences can influence visualization interpretation [29]. Larger and more diverse samples would help examine whether explanation preferences vary across user groups. Second, because the evaluation focused on subjective ratings and rankings, the findings speak to perceived explanation quality rather than analytical accuracy or efficiency. Third, the study examined two document collections, and additional evaluations would be needed to assess broader domain generality. Finally, CAPE explains layouts as presented and relies on heuristic pattern identification, so explanations may inherit biases or distortions from the underlying layout and may not capture all meaningful structures in more complex settings.

9 Conclusion

We presented CAPE, a context-aware explanation framework for spatialized document layouts that supports interpretation of spatial organization rather than explanation of how layouts are generated. By integrating semantic information with layout-derived spatial context, CAPE generates natural-language explanations that help users understand how documents and regions relate within a layout.

A controlled user study suggests that spatially grounded explanations are perceived as more helpful for understanding the layout compared with content-only baselines. These findings highlight the importance of incorporating spatial context into explanation mechanisms for spatial representations and suggest opportunities for supporting interpretation in broader exploratory analysis workflows.

Use of Generative AI

During the preparation of this work, the authors used ChatGPT to improve writing. After using this tool, the authors reviewed and edited the content as needed and take full responsibility for the content of the publication.

References

- [1] Christopher Andrews, Alex Endert, and Chris North. 2010. Space to think: large high-resolution displays for sensemaking. In *Proceedings of the SIGCHI Conference on Human Factors in Computing Systems* (Atlanta, Georgia, USA) (CHI '10). Association for Computing Machinery, New York, NY, USA, 55–64. doi:10.1145/1753326.1753336
- [2] Christopher Andrews and Chris North. 2013. The Impact of Physical Navigation on Spatial Organization for Sensemaking. *IEEE Transactions on Visualization and Computer Graphics* 19, 12 (2013), 2207–2216. doi:10.1109/TVCG.2013.205
- [3] Alexander Bendeck and John Stasko. 2025. An Empirical Evaluation of the GPT-4 Multimodal Language Model on Visualization Literacy Tasks. *IEEE Transactions on Visualization and Computer Graphics* 31, 1 (2025), 1105–1115. doi:10.1109/TVCG.2024.3456155
- [4] Yali Bian and Chris North. 2021. DeepSI: Interactive Deep Learning for Semantic Interaction. In *Proceedings of the 26th International Conference on Intelligent User Interfaces* (College Station, TX, USA) (IUI '21). Association for Computing Machinery, New York, NY, USA, 197–207. doi:10.1145/3397481.3450670
- [5] Yali Bian, Chris North, Eric Krokos, and Sarah Joseph. 2021. Semantic Explanation of Interactive Dimensionality Reduction. In *2021 IEEE Visualization Conference (VIS)*. 26–30. doi:10.1109/VIS49827.2021.9623322
- [6] Raphael Buchmüller, Dennis Collaris, Linhao Meng, and Angelos Chatzimparmpas. 2026. LangLasso: Interactive Cluster Descriptions through LLM Explanation. arXiv:2601.10458 [cs.HC] <https://arxiv.org/abs/2601.10458>
- [7] Kiroong Choe, Chaerin Lee, Soohyun Lee, Jiwon Song, Aeri Cho, Nam Wook Kim, and Jinwook Seo. 2025. Enhancing Data Literacy On-Demand: LLMs as Guides for Novices in Chart Interpretation. *IEEE Transactions on Visualization and Computer Graphics* 31, 9 (2025), 4712–4727. doi:10.1109/TVCG.2024.3413195
- [8] Michelle Dowling, Nathan Wycoff, Brian Mayer, John Wenskovitch, Scotland Leman, Leanna House, Nicholas Polys, Chris North, and Peter Hauck. 2019. Interactive visual analytics for sensemaking with big text. *Big Data Research* 16 (2019), 49–58. doi:10.1016/j.bdr.2019.04.003
- [9] Alex Endert, Russ Burtner, Nick Cramer, Ralph Perko, Shawn Hampton, and Kristin Cook. 2013. Typograph: Multiscale spatial exploration of text documents. In *2013 IEEE International Conference on Big Data*. 17–24. doi:10.1109/BigData.2013.6691709
- [10] Alex Endert, Patrick Fiaux, and Chris North. 2012. Semantic Interaction for Sensemaking: Inferring Analytical Reasoning for Model Steering. *IEEE Transactions on Visualization and Computer Graphics* 18, 12 (2012), 2879–2888. doi:10.1109/TVCG.2012.260
- [11] Alex Endert, Patrick Fiaux, and Chris North. 2012. Semantic interaction for visual text analytics. In *Proceedings of the SIGCHI Conference on Human Factors in Computing Systems* (Austin, Texas, USA) (CHI '12). Association for Computing Machinery, New York, NY, USA, 473–482. doi:10.1145/2207676.2207741
- [12] Alex Endert, Seth Fox, Dipayan Maiti, Scotland Leman, and Chris North. 2012. The semantics of clustering: analysis of user-generated spatializations of text documents. In *Proceedings of the International Working Conference on Advanced Visual Interfaces* (Capri Island, Italy) (AVI '12). Association for Computing Machinery, New York, NY, USA, 555–562. doi:10.1145/2254556.2254660
- [13] Rebecca Faust, David Glickenstein, and Carlos Scheidegger. 2019. DimReader: Axis lines that explain non-linear projections. *IEEE Transactions on Visualization and Computer Graphics* 25, 1 (2019), 481–490. doi:10.1109/TVCG.2018.2865194
- [14] Rao Fu, Jingyu Liu, Xilun Chen, Yixin Nie, and Wenhan Xiong. 2024. Scene-LLM: Extending Language Model for 3D Visual Understanding and Reasoning. arXiv:2403.11401 [cs.CV] <https://arxiv.org/abs/2403.11401>
- [15] Fausto German, Brian Keith, and Chris North. 2025. Narrative Trails: A Method for Coherent Storyline Extraction via Maximum Capacity Path Optimization. arXiv:2503.15681 [cs.IR] <https://arxiv.org/abs/2503.15681>

- [16] Carsten Görg, Zhicheng Liu, Jaeyeon Kihm, Jaegul Choo, Haesun Park, and John Stasko. 2013. Combining Computational Analyses and Interactive Visualization for Document Exploration and Sensemaking in Jigsaw. *IEEE Transactions on Visualization and Computer Graphics* 19, 10 (2013), 1646–1663. doi:10.1109/TVCG.2012.324
- [17] Jeffrey Heer and Ben Shneiderman. 2012. Interactive dynamics for visual analysis: A taxonomy of tools that support the fluent and flexible use of visualizations. *Queue* 10, 2 (2012), 30–55.
- [18] Florian Heimerl, Steffen Lohmann, Simon Lange, and Thomas Ertl. 2014. Word Cloud Explorer: Text Analytics Based on Word Clouds. In *2014 47th Hawaii International Conference on System Sciences*. 1833–1842. doi:10.1109/HICSS.2014.231
- [19] Z. Huang, D. Witschard, K. Kucher, and A. Kerren. 2023. VA + Embeddings STAR: A State-of-the-Art Report on the Use of Embeddings in Visual Analytics. *Computer Graphics Forum* 42, 3 (2023), 539–571. doi:10.1111/cgf.14859
- [20] IEEE VIS. 2026. IEEE Visualization and Visual Analytics Conference. <https://ieeevis.org/> Retrieved May 10, 2026.
- [21] Hyeon Jeon, Jeongin Park, Sungbok Shin, and Jinwook Seo. 2025. Stop Misusing t-SNE and UMAP for Visual Analytics. arXiv:2506.08725 [cs.HC] <https://arxiv.org/abs/2506.08725>
- [22] Daniel Keim, Gennady Andrienko, Jean-Daniel Fekete, Carsten Görg, Jörn Kohlhammer, and Guy Melançon. 2008. *Visual Analytics: Definition, Process, and Challenges*. Springer Berlin Heidelberg, Berlin, Heidelberg, 154–175. doi:10.1007/978-3-540-70956-5_7
- [23] Jason S. Kessler. 2017. Scattertext: a Browser-Based Tool for Visualizing how Corpora Differ. arXiv:1703.00565 [cs.CL] <https://arxiv.org/abs/1703.00565>
- [24] Ken Lang. 1995. Newsweeder: Learning to filter netnews. Available at <http://people.csail.mit.edu/jrennie/20NewsGroups/>.
- [25] Haitao Li, Qian Dong, Junjie Chen, Huixue Su, Yujia Zhou, Qingyao Ai, Ziyi Ye, and Yiqun Liu. 2024. LLMs-as-Judges: A Comprehensive Survey on LLM-based Evaluation Methods. arXiv:2412.05579 [cs.CL] <https://arxiv.org/abs/2412.05579>
- [26] Lee Lisle, Kylie Davidson, Edward J.K. Gitre, Chris North, and Doug A. Bowman. 2021. Sensemaking Strategies with Immersive Space to Think. In *2021 IEEE Virtual Reality and 3D User Interfaces (VR)*. 529–537. doi:10.1109/VR50410.2021.00077
- [27] Wei Liu, Eric Krokos, Kirsten Whitley, Rebecca Faust, and Chris North. 2026. LLM-Augmented Semantic Steering of Text Embedding Projection Spaces. arXiv preprint arXiv:2605.01957 (2026).
- [28] Wei Liu, Chris North, and Rebecca Faust. 2024. Visualizing Spatial Semantics of Dimensionally Reduced Text Embeddings. arXiv:2409.03949 [cs.HC] <https://arxiv.org/abs/2409.03949>
- [29] Zhengliang Liu, R. Jordan Crouser, and Alvitta Ottley. 2020. Survey on Individual Differences in Visualization. *Computer Graphics Forum* 39, 3 (2020), 693–712. doi:10.1111/cgf.14033
- [30] Scott M Lundberg and Su-In Lee. 2017. A Unified Approach to Interpreting Model Predictions. In *Advances in Neural Information Processing Systems*, I. Guyon, U. Von Luxburg, S. Bengio, H. Wallach, R. Fergus, S. Vishwanathan, and R. Garnett (Eds.), Vol. 30. Curran Associates, Inc. https://proceedings.neurips.cc/paper_files/paper/2017/file/8a20a8621978632d76c43dfd28b67767-Paper.pdf
- [31] Leland McInnes, John Healy, and James Melville. 2018. UMAP: Uniform Manifold Approximation and Projection for Dimension Reduction. arXiv:1802.03426 [stat.ML] <https://arxiv.org/abs/1802.03426>
- [32] Lan Huong Nguyen and Susan Holmes. 2019. Ten quick tips for effective dimensionality reduction. *PLOS Computational Biology* 15, 6 (06 2019), 1–19. doi:10.1371/journal.pcbi.1006907
- [33] Kwanrutai Nokkaew and Rachada Kongkachandra. 2018. Keyphrase Extraction as Topic Identification Using Term Frequency and Synonymous Term Grouping. In *2018 International Joint Symposium on Artificial Intelligence and Natural Language Processing (ISAI-NLP)*. 1–6. doi:10.1109/ISAI-NLP.2018.8693001
- [34] OpenAI. 2026. OpenAI API Documentation. <https://platform.openai.com/docs> Retrieved May 10, 2026.
- [35] Rajvardhan Patil, Sorio Boit, Venkat Gudivada, and Jagadeesh Nandigam. 2023. A Survey of Text Representation and Embedding Techniques in NLP. *IEEE Access* 11 (2023), 36120–36146. doi:10.1109/ACCESS.2023.3266377
- [36] Robert Pienta, James Abello, Minsuk Kahng, and Duen Horng Chau. 2015. Scalable graph exploration and visualization: Sensemaking challenges and opportunities. In *2015 International Conference on Big Data and Smart Computing (BIGCOMP)*. 271–278. doi:10.1109/35021BIGCOMP.2015.7072812
- [37] Peter Pirolli and Stuart Card. 2005. The sensemaking process and leverage points for analyst technology as identified through cognitive task analysis. In *Proceedings of international conference on intelligence analysis*, Vol. 5. McLean, VA, USA, 2–4.
- [38] Shivam Raval, Carolyn Wang, Fernanda Viégas, and Martin Wattenberg. 2023. Explain-and-Test: An Interactive Machine Learning Framework for Exploring Text Embeddings. In *2023 IEEE Visualization and Visual Analytics (VIS)*. 216–220. doi:10.1109/VIS54172.2023.00052
- [39] Marco Tulio Ribeiro, Sameer Singh, and Carlos Guestrin. 2016. “Why Should I Trust You?”: Explaining the Predictions of Any Classifier. In *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining* (San Francisco, California, USA) (*KDD '16*). Association for Computing Machinery, New York, NY, USA, 1135–1144. doi:10.1145/2939672.2939778
- [40] Tuukka Ruotsalo, Jaakko Peltonen, Manuel Eugster, Dorota Glowacka, Ksenia Konyushkova, Kumaripaba Athukorala, Ilkka Kosunen, Aki Reijonen, Petri Myllymäki, Giulio Jacucci, and Samuel Kaski. 2013. Directing exploratory search with interactive intent modeling. In *Proceedings of the 22nd ACM International Conference on Information & Knowledge Management* (San Francisco, California, USA) (*CIKM '13*). Association for Computing Machinery, New York, NY, USA, 1759–1764. doi:10.1145/2505515.2505644
- [41] Gerard Salton and Christopher Buckley. 1988. Term-weighting approaches in automatic text retrieval. *Information Processing & Management* 24, 5 (1988), 513–523. doi:10.1016/0306-4573(88)90021-0
- [42] Ksh Singh, H Mamata Devi, and Anjana Kakoti Mahanta. 2017. Document representation techniques and their effect on the document Clustering and Classification: A Review. *International Journal of Advanced Research in Computer Science* 8, 5 (2017).
- [43] Daniel Smilkov, Nikhil Thorat, Charles Nicholson, Emily Reif, Fernanda B. Viégas, and Martin Wattenberg. 2016. Embedding Projector: Interactive Visualization and Interpretation of Embeddings. arXiv:1611.05469 [stat.ML] <https://arxiv.org/abs/1611.05469>
- [44] Sangho Suh, Bryan Min, Srishti Palani, and Haijun Xia. 2023. Sensecape: Enabling Multilevel Exploration and Sensemaking with Large Language Models. In *Proceedings of the 36th Annual ACM Symposium on User Interface Software and Technology* (San Francisco, CA, USA) (*UIST '23*). Association for Computing Machinery, New York, NY, USA, Article 1, 18 pages. doi:10.1145/3586183.3606756
- [45] Yuan Sui, Mengyu Zhou, Mingjie Zhou, Shi Han, and Dongmei Zhang. 2024. Table Meets LLM: Can Large Language Models Understand Structured Table Data? A Benchmark and Empirical Study. In *Proceedings of the 17th ACM International Conference on Web Search and Data Mining* (Merida, Mexico) (*WSDM '24*). Association for Computing Machinery, New York, NY, USA, 645–654. doi:10.1145/3616855.3635752
- [46] Flora S. Tsai. 2011. Dimensionality reduction techniques for blog visualization. *Expert Systems with Applications* 38, 3 (2011), 2766–2773. doi:10.1016/j.eswa.2010.08.067
- [47] Laurens Van der Maaten and Geoffrey Hinton. 2008. Visualizing data using t-SNE. *Journal of machine learning research* 9, 11 (2008).
- [48] Xiting Wang, Shixia Liu, Junlin Liu, Jianfei Chen, Jun Zhu, and Baining Guo. 2016. TopicPanorama: A Full Picture of Relevant Topics. *IEEE Transactions on Visualization and Computer Graphics* 22, 12 (2016), 2508–2521. doi:10.1109/TVCG.2016.2515592
- [49] James Wexler, Mahima Pushkarna, Tolga Bolukbasi, Martin Wattenberg, Fernanda Viégas, and Jimbo Wilson. 2020. The What-If Tool: Interactive Probing of Machine Learning Models. *IEEE Transactions on Visualization and Computer Graphics* 26, 1 (2020), 56–65. doi:10.1109/TVCG.2019.2934619
- [50] J.A. Wise, J.J. Thomas, K. Pennock, D. Lantrip, M. Pottier, A. Schur, and V. Crow. 1995. Visualizing the non-visual: spatial analysis and interaction with information from text documents. In *Proceedings of Visualization 1995 Conference*. 51–58. doi:10.1109/INFVIS.1995.528686
- [51] Yang Zhang, Hanlei Jin, Dan Meng, Jun Wang, and Jinghua Tan. 2025. A Comprehensive Survey on Process-Oriented Automatic Text Summarization with Exploration of LLM-Based Methods. arXiv:2403.02901 [cs.AI] <https://arxiv.org/abs/2403.02901>