

Brushing and Interpretive Linking for Knowledge Work

Zheng Ning

Microsoft Research
Redmond, WA, USA
University of Notre Dame
Notre Dame, IN, USA
zning@nd.edu

Michel Pahud

Microsoft Research
Redmond, WA, USA
mpahud@microsoft.com

Hugo Romat

Microsoft Research
Redmond, WA, USA
romathugo@microsoft.com

Nicolai Marquardt

Microsoft Research
Redmond, WA, USA
nicmarquardt@microsoft.com

Fanny Chevalier

Microsoft Research
Redmond, WA, USA
University of Toronto
Toronto, ON, Canada
fanny@dgp.toronto.edu

Jeevana Priya Inala

Microsoft Research
Redmond, WA, USA
jinala@microsoft.com

Chenglong Wang

Microsoft Research
Redmond, WA, USA
chenwang@microsoft.com

David Brown

Microsoft Research
Redmond, WA, USA
dabrown@microsoft.com

Toby Jia-Jun Li

University of Notre Dame
Notre Dame, IN, USA
toby.j.li@nd.edu

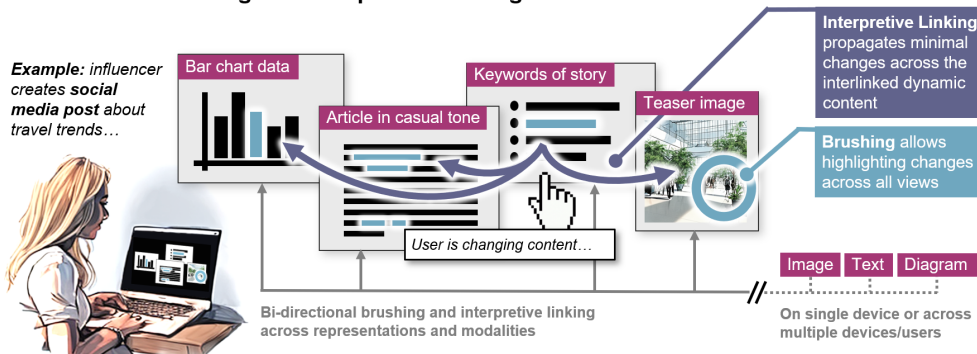
Ken Hinckley

Microsoft Research
Redmond, WA, USA
kenh@microsoft.com

Nathalie Riche

Microsoft Research
Redmond, WA, USA
nath@microsoft.com

Generalized Brushing and Interpretive Linking



Probe-to-Prototype Methodology

Eliciting concepts for brushing and interpretive linking from six professional content creators...

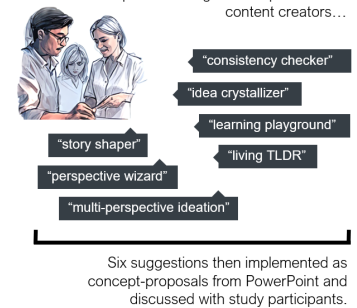


Figure 1: (left) Overview of brushing and interpretive linking across representations and modalities (e.g., text, images, charts), leveraging generative AI to establish conceptual relationships and propagating changes; (right) Examples of workflow concepts elicited from professional knowledge workers, illustrating how the approach generalizes across diverse use cases.

Abstract

This work builds on brushing and linking from data visualization and extends it to knowledge work. Knowledge workers frequently repurpose content across artifacts (e.g., creating a one-page summary of a report and visuals for its presentation), yet the conceptual relationships between these representations are often implicit, placing a cognitive burden on knowledge workers to maintain alignment. We introduce *interpretive linking*, an approach that

leverages generative AI to identify and highlight semantically related content across representations (*brushing*) and to interpret and propagate changes between them (*linking*). Rather than preserving strict correspondence, interpretive linking aims to maintain semantic coherence across formats and modalities. Using a research-through-design approach, we built a technology probe leveraging multiple models across textual and visual content and experimented with different strategies for managing semantic coherence. We then conducted a qualitative study with six knowledge workers, who used the probe to reflect on potential workflows. Participants generated 22 application concepts, ranging from responsive documents adapting to device form factor to collaborative tools supporting parallel work across modalities. Together, these findings surface design opportunities and open challenges for interpretive linking in knowledge work.



This work is licensed under a Creative Commons Attribution 4.0 International License.
Graphics Interface 2026 (GI'26), Waterloo, Canada
© 2026 Copyright held by the owner/author(s).
ACM ISBN 978-1-4503-9157-0/00/00
<https://doi.org/10.1145/3491102.0000000>

CCS Concepts

• **Human-centered computing** → **Interaction paradigms**.

Keywords

content transformation, generative AI, knowledge work

ACM Reference Format:

Zheng Ning, Michel Pahud, Hugo Romat, Nicolai Marquardt, Fanny Chevalier, Jeevana Priya Inala, Chenglong Wang, David Brown, Toby Jia-Jun Li, Ken Hinckley, and Nathalie Riche. 2026. Brushing and Interpretive Linking for Knowledge Work. In *Graphics Interface 2026 (GI'26)*. ACM, New York, NY, USA, 11 pages. <https://doi.org/10.1145/3491102.0000000>

1 Introduction

Brushing and linking is a foundational technique in data visualization that enables analysts to explore data across multiple coordinated views [1, 23]. Brushing allows a user to select or isolate data within one view. Linking ensures that these highlighting and filtering interactions are dynamically reflected in all the other views. This strategy allows analysts to explore data from multiple perspectives, leading to deeper insights and a more comprehensive understanding [26].

Knowledge workers frequently repurpose material, continually adapting content to suit different audiences, media, and collaborative contexts [19]. For example, they may turn a long document into a summary, a representative diagram, and a slide deck [11, 16]. When repurposing material, they perform a mental brushing and linking concurrently with content editing, such as reading a section of the report to write corresponding summary lines or slides. Unlike data visualization, where data is indexed with consistent identifiers, supporting brushing and linking across textual and visual content is conceptually and technically more challenging because of the inherent ambiguities in what constitutes ‘the same’ content across artifacts that differ in modality, granularity, and overall communicative conventions.

In this paper, we introduce *interpretive linking*, a conceptual approach that reframes linking from preserving exact correspondence to managing semantic coherence across representations. In particular, we explore how generative models can support *brushing* (identifying and highlighting semantically related content across representations) and *linking* (interpreting and propagating changes between those representations). Throughout, we use *interpretive linking* to refer to this end-to-end flow that encompasses both mechanisms. We explore interpretive linking in the context of textual and visual content, but as a general approach, it is applicable across a range of knowledge work scenarios with different modalities and formats. For integration of interpretive linking in these scenarios, we propose a set of strategies to preserve meaning and communicative intent during change propagation.

Adopting a research-through-design approach [37], we built a technology probe as a vehicle for exploring how interpretive linking can be realized and appropriated in practice. We then conducted a qualitative study with six knowledge workers, who engaged with the probe to reflect on how interpretive linking could reshape their workflows or enable new scenarios. Participants generated 22 application concepts, ranging from responsive documents that maintain coherence across representations to collaborative ideation tools that

support parallel work across modalities. Insights from our probe iterations and participant engagements surface design opportunities and open challenges for interpretive linking in knowledge work.

In summary, we contribute:

- **The concept of *interpretive linking***: a reframing of brushing and linking from preserving exact correspondence to managing semantic coherence and communicative intent across representations, mediated by generative AI.
- **A technology probe and a set of design strategies** (meta-descriptions as soft constraints, version-aware propagation, and visual traceability of changes) that demonstrate how interpretive linking can be realized across textual and visual content.
- **Exploratory design insights from a probe-to-prototype study with six professional content creators**, including 22 participant-generated workflow concepts and reflections that surface opportunities, open challenges, and tensions around control, trust, and propagation.

2 Related Work

To situate our work, we review three strands of research. First, we revisit brushing and linking, a foundational technique in visualization that coordinates data across multiple views (Sec. 2.1). We then review knowledge workers’ practices and the opportunities this creates for extending brushing and linking to this domain (Sec. 2.2). Finally, we discuss recent technical advancements in content transformation and multi-modal content generation that form the foundation for realizing brushing and interpretive linking in knowledge work (Sec. 2.3).

2.1 Brushing & Linking in Visualization

Brushing and linking is a foundational interaction technique in the visualization community [23]. Brushing refers to selecting or highlighting a subset of data in one representation, while linking ensures that the corresponding data are simultaneously revealed or filtered in other representations. When applied across multiple coordinated views, brushing and linking allow analysts to fluidly trace relationships, compare perspectives, and propagate selections across heterogeneous data displays [5, 23]. This mechanism has been shown to greatly enhance exploratory analysis by reducing the cognitive burden of manual cross-referencing and by supporting richer sense-making processes. Numerous systems exemplify its value: image editing tools that couple histograms and pixel views for interactive adjustment [7, 10], or complex visualization environments that integrate scatterplots with embodied navigation through high-dimensional data spaces [9]. Our work focuses on exploring the feasibility and potential use cases of brushing and linking in knowledge work, across textual and visual content.

2.2 Repurposing in Knowledge Work

Similar to data analysts who work across multiple coordinated views when exploring datasets, knowledge workers routinely juggle diverse representations of the same underlying material. For example, writers reworking long-form text into concise summaries [35], journalists translating print articles into multi-modal video scripts and storyboards [34], programmers shifting between sketches, comments, and executable code [36], or converting raw data into charts

[33]. While these transformations are essential, they typically require repeated manual effort to coordinate different pieces of content.

In this work, we use repurposing to refer broadly to these transformations, encompassing the ongoing adaptation, reframing, and coordination of content across representations, modalities, and contexts. The content being transformed spans modalities (e.g., text, images) and formats within a modality (e.g., paragraph, bullet points). Workers engage in these transformations with different goals. From a single-user workflow perspective, transformations help fluidly move between granular and abstracted views, improving efficiency and recall through richer mental representations [8, 20]. They also reduce cognitive load by distributing information across complementary channels [29]. From a collaborative perspective, transformations allow multiple users to work with representations that suit their distinct roles, preferences, or cognitive styles [18, 25, 30]. More recent design research has explored how multimodal input can better align professional creators' intent with generative-AI output, such as in DesignPrompt [22] which combines image, color palettes, and text in unified prompts. Additionally, from a device and environment perspective, knowledge workers frequently adapt content to the affordances of different form factors and contexts, such as tablets for sketching, or wearables for quick access, thereby extending interaction seamlessly across settings and supporting situated cognition [3, 13, 14].

Recognizing the widespread practice of content transformation in knowledge work and the effectiveness of brushing and linking in visualization, we are motivated to ask whether similar strategies could extend to knowledge work. In this research, we explored how such strategies can be realized by leveraging generative AI to identify related content in different representations and propagate changes while preserving core meaning and communicative intent of each representation.

2.3 Content Transformation Pipelines

Earlier approaches to content transformation assumed explicit, pre-defined mappings between representations, exemplified by optical character recognition, text-to-speech and speech-to-text systems, and rule-based format transformations. Within HCI, researchers extended these approaches through interactive systems that supported repurposing across applications, devices, and collaborative settings. Examples include sketch-to-UI tools that translated wireframes into executable interfaces [27], “fluid documents” that enabled styling and co-editing [6], and cross-device platforms that supported moving and adapting content to different form factors and collaborative settings [3, 13, 18].

More recently, advances in large pretrained generative models have expanded the range of possible transformations, including text-to-text and text-to-code via large language models [17, 32], text-to-image synthesis via diffusion models [24], and emerging approaches for text-to-audio and text-to-video [12, 31]. These models demonstrate strong capabilities for synthesizing and repurposing content across modalities. Researchers have also begun structuring LLM interactions to balance exploration with control, for example by defining explicit semantic dimensions that shape and diversify generation [28]. However, much of this work has focused on isolated

transformations, leaving open questions about how such transformations can be coordinated, interpreted, and maintained as content evolves across multiple representations in real knowledge-work practices. This gap shifts attention from what transformations are possible to how they are structured and coordinated in use. Recent work by Cao et al. [4] addresses this directly, by proposing compositional structures as substrates that scaffold human-AI co-creation across linked views.

In this paper, we explore two fundamental aspects of brushing and linking that generative models now make possible in textual and visual content: identifying semantically related concepts across evolving artifacts, and interpreting and propagating change while preserving semantic coherence and communicative intent across representations.

3 Interpretive Linking

Linking refers to the propagation of interactions or changes across multiple representations such that actions performed in one representation are reflected in others. In information visualization, linking is commonly used to synchronize state across coordinated views, enabling integrated exploration of related representations through well-defined mappings. Extending linking to knowledge work poses two core challenges.

Mapping. The first challenge concerns identifying which elements across representations should be linked in the first place. In visualization systems, this identification is typically straightforward: marks correspond to shared data items, and links are established through explicit identifiers or one-to-one mappings. In contrast, text and images rarely provide such stable structure. Relationships between passages, summaries, figures, or captions are often implicit, contextual, and many-to-many, rather than defined by fixed identifiers.

Generative AI models support establishing these relationships by identifying semantically related concepts across representations, without extensive explicit specification. Large language models can relate a sentence in an abstract to corresponding explanatory passages, or connect a formally phrased legal statement with its plain-language explanation. Vision and image generation models can make similar connections across textual and visual content, such as linking regions of an image to captions, discussions, or explanatory descriptions of charts and diagrams.

Propagation. Whereas mapping concerns identifying what should be linked, propagation concerns what should change once those links are in place. In information visualization, linking primarily propagates selection or filtering state: highlighting a subset of data in one view reveals or suppresses corresponding items in other views, while the underlying data remains unchanged. This form of propagation supports coordinated exploration, but is largely limited to synchronizing views rather than modifying content.

In knowledge work, however, propagation often involves changes to content rather than to view state. Authors revise text, refine arguments, and update visuals as ideas evolve. Extending linking to these settings raises questions about how edits in one representation should be interpreted and reflected elsewhere while preserving semantic coherence. For example, revising a concept in one section of

a document could prompt updates in related passages, or adjusting a figure could require corresponding changes to captions and references. Across documents, linking could help keep content aligned across languages or registers, such as maintaining consistency between a technical report and a legally phrased patent. Supporting such propagation enables coordinated, cross-modal authoring and sensemaking beyond what traditional linking affords.

Interpretive linking. We refer to this form of generative-AI-enabled linking as *interpretive linking*. In interpretive linking, both the identification of related elements and the propagation of change are inherently fluid and ambiguous. Relationships across representations are inferred rather than fixed, and changes do not follow pre-defined correspondences. This ambiguity is further compounded by the non-deterministic behavior of generative AI models, which may yield different or partially incorrect interpretations across iterations. As a result, interpretive linking introduces risks such as semantic drift or inadvertently introducing hallucinations. At the same time, it enables forms of cross-modal and cross-register editing (i.e., across different communicative roles or styles) that are not possible with traditional, programmed linking. Interpretive linking therefore shifts the problem from preserving exact correspondence to developing strategies that maintain semantic coherence and the communicative intent of each representation as content evolves and changes are propagated across artifacts.

As a conceptual approach, interpretive linking is not bound to any specific model architecture or pipeline. It requires (1) a mechanism for inferring semantic relationships across diverse representations, and (2) means of propagating changes across the artifacts. Bidirectionality, while valuable in several scenarios we describe in this paper, is not a strict requirement, as interpretive linking can also operate unidirectionally. Similarly, the approach is agnostic to the specific content modalities used, and extends in principle to audio, video, code, and other artifact representations beyond text and images.

Extended Brushing. Whereas the mapping and propagation challenges concern what the system infers and what it changes, *brushing* concerns how the user invokes the mapping in the first place. Adapted from information visualization [1, 23], brushing in knowledge work allows authors to select a piece of content in one representation and surface its semantically related counterparts in the others, without committing to any edit. For example, brushing a sentence in an abstract can reveal the paragraphs that develop the idea later in the document; brushing a claim in a formally phrased patent can surface the corresponding plain-language passages in an explanatory report. Brushing can also operate across modalities, where a brushed region of an image can highlight the text passages that describe or discuss it, while a brushed sentence can surface the figures, charts, or visual examples that concretize the selected concept. By making the AI-inferred semantic relationships visible without triggering any modification, brushing supports non-linear reading, comparison across formats, and sensemaking as authors move between representations. It thus serves as the counterpart to propagation within interpretive linking: propagation commits to changes, brushing reveals relationships.

4 Technology Probe

To explore interpretive linking in practice, given its fluid and non-deterministic nature, we adopted a research-through-design approach [37] and built a technology probe as an experimental artifact. This technology probe is a vehicle for investigating how interpretive linking can be realized across textual and visual content, rather than a system intended for deployment. Through iterative development, we refined both the model pipelines and interaction strategies embodied in the probe, using it to investigate how interpretive linking might operate across text and images, particularly with respect to maintaining semantic coherence.

The probe also enabled us to observe how knowledge workers appropriate these capabilities in their own scenarios, revealing workflows, breakdowns, and opportunities for brushing and interpretive linking that extend beyond cases we could anticipate through design alone.

4.1 Probe Design Overview

We developed a web-based, multi-user technology probe that enables authors to create, edit, and coordinate multiple pieces of content with interpretive linking on an infinite canvas (Figure 2). The fundamental unit of interaction is a content block, in which users specify both the modality (text, chart, or image) and a free-form meta-description of the content (e.g., “bullet points,” “bar chart,” “cover for children’s book”). Users can create, resize, and spatially arrange multiple content blocks, which are all linked upon creation.

The canvas functions as a desktop-like environment composed of multiple simultaneous views, where each content block represents a distinct artifact with its own communicative intent and role (e.g., summary, explanation, visual illustration). This design enables users to repurpose shared content into multiple representations while continuing to work with each artifact independently. In addition, users can quickly assemble, juxtapose, and manipulate related artifacts to explore diverse scenarios.

In practice, determining whether artifacts should be linked, temporarily decoupled, or connected through custom relationships is itself a complex design problem [21]. To isolate questions about interpretive linking and semantic coherence, the probe adopts a simplifying assumption that all content blocks on the canvas are connected. Exploring richer, user-controlled linking structures is beyond the scope of this probe.

4.2 Managing Semantic Coherence

A central challenge in interpretive linking is maintaining semantic coherence between representations as they evolve. The probe addresses this challenge through several design strategies.

Meta-Descriptions. Associated with each content block, meta-descriptions act as soft constraints that guide how content is translated, helping preserve the communicative intent and format of each representation. In the probe, these meta-descriptions are explicitly specified by users or are generated the first time the user adds content. Rather than defining extensive transformation rules between artifacts, meta-descriptions capture high-level characterizations of an artifact, possibly tone, level of abstraction, or visual style, which the system attempts to preserve during transformation.

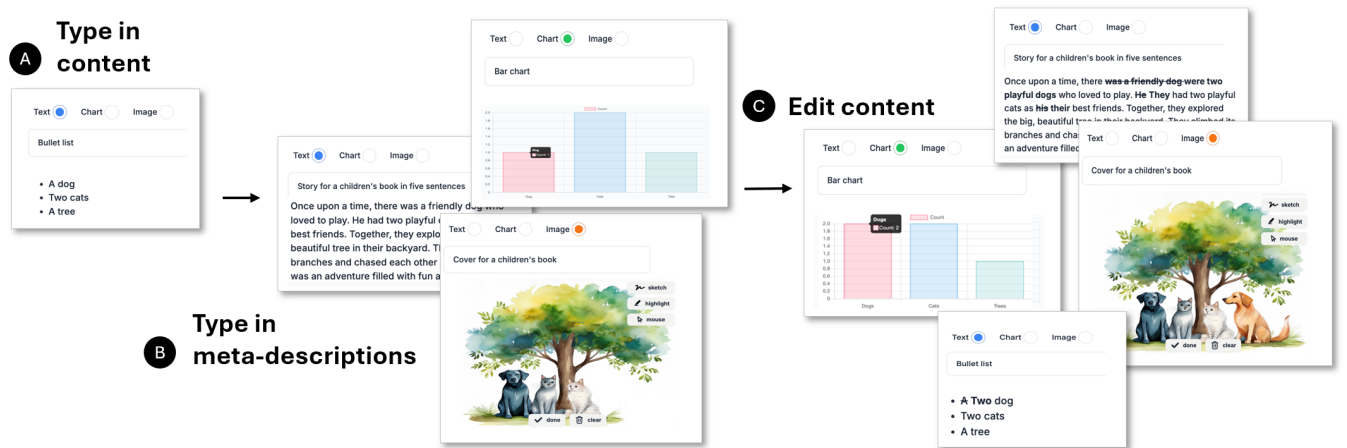


Figure 2: Example usage of the probe. (A) User creates a text block and types in content, meta-description is updated. (B) User creates three more blocks and enters meta-descriptions; their content is generated by leveraging meta-descriptions and content of the initial block. (C) User edits a value in the chart directly, changes are propagated to the other content blocks in the canvas.

This expands on related approaches that structure LLM generation via explicit semantic dimensions [28], where we leverage such dimensions as soft constraints on cross-modal propagation.

For example, a text block described as “bullet points for a presentation” or an image block described as a “cover illustration” establishes expectations about structure and emphasis that differ from those of a narrative paragraph or an analytical chart. During linking, these meta-descriptions bias transformations toward outputs that remain appropriate for their representational context, even as underlying content changes. Treating meta-descriptions as soft and interpretable constraints allows the system to preserve communicative intent without rigid coupling, while supporting workflows in which representations are incrementally derived, repurposed, and refined.

Versioning. To support minimal yet meaningful change, the transformation process incorporates not only the current source content and the target’s meta-description, but also the previous versions of both source and target content. This historical context allows the system to infer which aspects of the representation should be preserved and which should be updated. For example, when a numerical value is modified in a chart, only the semantically corresponding portion of a linked text description is revised, while unrelated content remains unchanged.

Visual Traceability of Change. In addition to shaping how changes are generated, the probe makes transformations visually traceable so that users can track how representations evolve over time. When linking propagates a change, updates are surfaced locally. For text, modified segments are marked in bold or striked through; for images and charts, corresponding regions are highlighted. These visual cues help users identify what has changed and how a transformation relates to the original content.

Brushing further supports traceability by allowing users to explicitly reveal linked content across representations. By selecting

text, chart elements, or image regions, users can surface related elements in other blocks, helping them maintain mental connections between representations.

4.3 Interpretive Linking Pipelines

Interpretive linking is bidirectional and supports transformations across modalities. Changes in text can be reflected in charts or images, and vice versa, using combinations of large language models, segmentation models, and diffusion-based image generation models. Figure 3 illustrates the process of propagating a change made to a content block to a linked content block in the canvas.

Text-text (and text-chart). Meta-descriptions associated with each content block are incorporated directly into model calls and act as soft constraints that help preserve key communicative characteristics during linking. These meta-descriptions may specify textual structure (e.g., “bullet points”), visual style (e.g., “watercolor children’s illustration”), or format (e.g., “horizontal bar chart”), allowing linking prompts to explicitly maintain these properties and limit semantic or stylistic drift. To further support minimal yet meaningful change propagation, the pipeline also provides prior versions of both linked blocks to the language model, leveraging generative models’ ability to reason about similarities and differences across artifacts. This history enables the model to infer additional constraints that may not be explicitly captured in meta-descriptions, while keeping user-provided descriptions succinct and high-level. Concretely, each linking operation invokes an LLM with the content of both blocks, the meta-description of the target block as a soft constraint on its communicative style, and the version history to preserve semantic coherence while propagating changes.

Linking between a text block and a chart block follows the same approach, with the difference being that we prompt the LLM to generate HTML code for the chart or diagram.

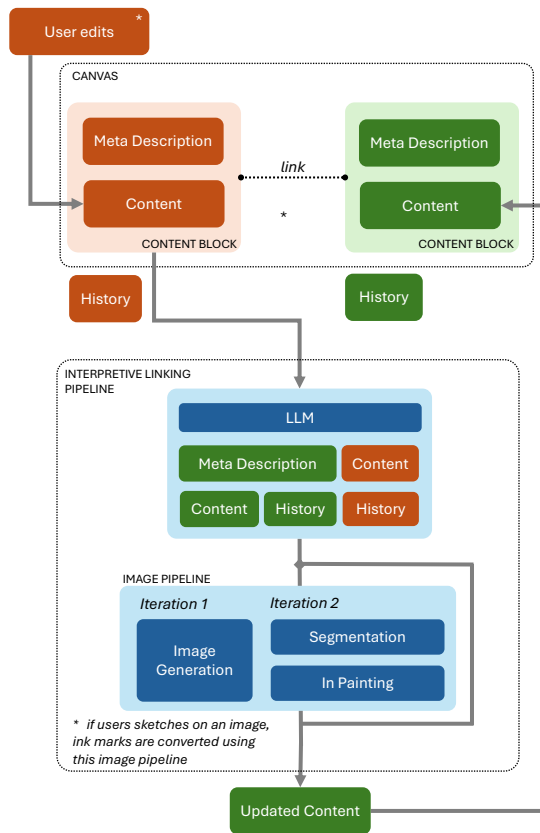


Figure 3: Interpretive Linking Pipeline. User edits content of the red block. Changes are propagated to linked green block. The pipeline consists of a call to an LLM, and, if the green block is an image, the LLM generates calls to image models.

Text - Image. Cross-modality text-image interpretive linking involves invoking additional image models, with prompts generated by the LLM. Through iterative development of the probe, we explored two pipeline strategies. An initial approach based on direct image generation frequently produced subtle but unintended visual changes, despite efforts to localize edits through prompt engineering. To address this, we explored a more controlled pipeline that combined image segmentation with inpainting using a Stable Diffusion model, allowing changes to be spatially constrained while preserving surrounding visual context (Figure 4).

In this second iteration, when a change originates from a text block (e.g., replacing “cat” with “elephant”), the system first analyzes the source image to capture its overall visual style. The LLM then generates a prompt that specifies the desired change while preserving this style. A segmentation model identifies the object referenced in the text and produces a corresponding mask, after which a diffusion model inpaints the new object within the masked region. This approach constrains changes spatially while maintaining consistency with the surrounding visual context.

In the reverse direction, when a user modifies an object in the image using a digital pen, the marked image is first processed to

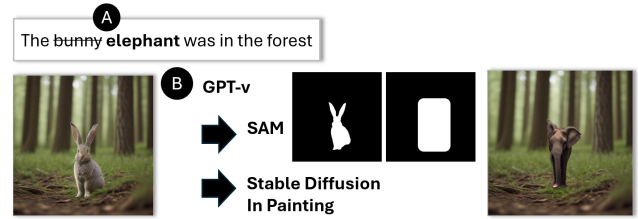


Figure 4: Propagating localized changes from text to image.

match the style of the original image. The updated visual content is then analyzed by a vision model, and the LLM updates the corresponding text description to reflect the new object while keeping the rest of the description unchanged.

Image - Image. To link between two visual blocks, we route edits through an LLM because the paired images may differ substantially in pose, viewpoint, or depiction, so the system must infer the semantics of the user’s edit (e.g., “add glasses to the cat”) rather than copy pixels; the LLM describes the intended change, which can then be localized and applied to the other image via segmentation and inpainting. When a user sketches on one image, the sketch is first transformed to match the style of the original image. The LLM then analyzes the modified image in the context of its source to infer the intended change and produces a textual description. Based on this description, a segmentation model identifies the corresponding region in the target image and generates a mask, after which a diffusion model performs inpainting within the masked region to render the new object while preserving the target image’s overall visual style and context.

4.4 Implementation

We developed this web-based technology probe using React¹ for the front-end and Node.js² for the back-end to synchronize the data across multiple users and devices. The computational pipeline utilizes GPT-4o³ and Stable Diffusion (v1.10.1). We used Azure OpenAI GPT-4o, model gpt-4o-2024-05-13, which was the latest version at the time of the experiment. The Stable Diffusion model runs on a local machine equipped with an NVIDIA GeForce RTX 4090 GPU and an Intel Core i9-14900K 3.20 GHz CPU. While we used these models to enable our technology probes, interpretive linking is not tied to a particular model or architecture.

5 User Study

To explore how brushing and interpretive linking could enable new forms of coordination and repurposing in knowledge work, we conducted a qualitative study with six professional content creators. Rather than evaluating a single predefined application, our goal was to use brushing and interpretive linking as a generative lens for uncovering emerging workflows, breakdowns, and opportunities across knowledge work practices.

¹<https://react.dev>

²<https://nodejs.org>

³<https://platform.openai.com/docs/models/gpt-4o>

5.1 Procedure

To gather a rich set of scenarios, we selected six participants (2M, 4F) from different backgrounds and roles (Table 1), a sample size more aligned with exploratory probe-based studies [15] aimed at breadth of ideation rather than statistical generalization.

Each participant took part in two online sessions conducted 5–7 days apart. We employed a *probe-to-prototype* approach, in which workflow ideas were captured as they emerged during a first ideation session (approx. 60 min), materialized by the research team using the probe between sessions, and then returned to participants for reflection in a second session (approx. 30 min).

Materializing abstract speculation into low-fidelity interactive prototypes enabled participants to more concretely assess opportunities, constraints, and breakdowns in use, as well as to provide informed feedback on the underlying generative AI pipeline by engaging directly with the probe and its behaviors in practice.

The first session (60 minutes) opened with a short video and live demonstration, followed by three example concept proposals the research team had built with the probe in advance. These examples were introduced to make the probe’s capabilities and constraints concrete without prescribing the space of ideas. Participants were then asked to propose workflow concepts and ideas grounded in their own practices. Concepts were captured as free-form descriptions and clarified through follow-up questions in the session. Across the six participants we collected a total of 22 distinct workflow concepts. Between sessions, the research team prototyped two of each participant’s concepts using the probe, selected to maximize diversity across coordination types (repurposing, reflection/ideation, collaboration). When specific functionalities or modalities (e.g., video) were not supported, we created the closest related variant. In the second session (30 minutes), participants experienced these materialized prototypes and reflected on design considerations, breakdowns, and opportunities.

5.2 Data Collection and Analysis

All sessions were conducted remotely via video call, recorded with participant consent, and transcribed. We performed a thematic analysis of the transcripts, using use cases, workflow concepts, benefits and design tension as deductive codes, while emerging themes were inductively captured. We focus our analysis on existence, following an interpretivist approach to reflexive thematic analysis [2].

5.3 Concept proposals

Across the study, participants proposed 22 workflow concepts that illustrate how interpretive linking can unlock new opportunities for knowledge work. We group the workflows below to highlight the kinds of new capabilities participants envisioned (Figure 5 shows representative examples).

Content repurposing. Participants’ discussions around derived artifacts reflected common repurposing practices described in prior work, where the same material is adapted across representations, audiences, and contexts. Several participants suggested that interpretive linking could help coordinate these related representations, which are otherwise maintained through manual rewriting. For example, P1 described transforming a research report into

presentation-oriented artifacts, while P3 and P4 discussed rewriting content such as feature descriptions or notifications from different perspectives (Figure 5 perspective wizard), including technical and non-technical audiences. Participants also noted that derived artifacts often serve different cognitive or communicative roles, such as learning a new programming language by experimenting side by side with a familiar one (Figure 5 learning playground). Together, these observations point to opportunities for systems that help knowledge workers manage and navigate multiple perspectives during repurposing work.

Beyond scenarios centered on content repurposing, participants surfaced additional insights about how interpretive linking might shape reflection, ideation, and collaboration in knowledge work.

Reflection and ideation. Several participants described workflows in which interpretive linking could help connect early, informal ideas with more structured outputs. For example, P1 and P3 discussed capturing brief notes on mobile devices during meetings or while away from their desks, and later transforming these fragments into slides or visual materials. Participants framed this as a way to keep early ideas connected to downstream artifacts as they evolve, rather than treating ideation and production as fully separate stages (Figure 5 idea crystallizer, story shaper).

Participants reflected that the brushing aspect of interpretive linking, which surfaces the same content across multiple representations, helped reframe their thinking, supporting reflection and perspective switching. P5 reflected that “*you could see many different ways of seeing the same information. That was really cool.*”, praising the value of seeing the same content in different forms for stimulating creativity. Similarly, P1 valued how seeing alternative representations made content feel like “*seeing text brought to life.*” especially when creating alternative representations would otherwise require substantial manual effort. P4 highlighted the ability to edit content bidirectionally across representations as particularly distinctive, describing this affordance as “*pretty unique . . . and super valuable.*” P6 further stressed the importance of grounding generated content in selected source materials and tracking provenance to maintain creative control and address copyright concerns.

Unlocking new collaborative dynamics. Participants also described collaborative scenarios, where interpretive linking might support coordination across roles and representations. P4 and P6 described maintaining shared outlines or indexes derived from artifacts owned by different team members (Figure 5 living TLDR). Others reflected on multi-perspective ideation, such as artists iterating on shared character concepts through different visual styles (P6), or interdisciplinary teams working across narrative text, storyboards, and scripts (P1).

Several participants framed interpretive linking as a mediating layer in collaboration: P3 described it as an “*AI collaborator.*” and P4 characterized it as a mediator that could “*transform content into something else.*” P1 and P5 found these collaborative scenarios particularly compelling, with P5 envisioning writers and designers working on the same content through different representations, describing this as “*a cool, different human–AI interaction paradigm.*” Together, these accounts suggest that interpretive linking enables collaborators to remain loosely aligned while contributing through

Participant	Role	Field of Expertise
P1	UX Researcher	conducts research on video games to inform the creative team
P2	Graphics Designer	design of intranet apps for employees of large organizations
P3	Content Designer	design of components in user interfaces such as menu labels, buttons
P4	Content Designer	design of user notifications within a suite of apps for information workers
P5	Project Manager	specification of goals and requirements for a team collaboration app
P6	Technical Artist	crafting images and programming 3D content for mixed reality experiences

Table 1: Background of our six study participants.

representations suited to their roles, expertise, and communicative needs, supporting coordination, reflection, and collaboration without enforcing uniform representations.

5.4 Reflections on interpretive linking

We synthesize below the reflections of participants after experiencing their own concepts materialized into interactive prototypes in the second session.

Cost-value tradeoff of propagating changes. Participants consistently framed the value of interpretive linking in relation to the effort and risk of coordinating consequential changes across artifacts. Interpretive linking was seen as most compelling when updates would otherwise be time-consuming or error-prone to perform manually. As P1 noted, *“Generative AI is better than me at making these kinds of changes—not small edits, but taking storyboards and turning them into a report. That would save hours.”* This framing positioned interpretive linking less as automating trivial edits and more as supporting changes with cascading effects.

Maintaining consistency across representations emerged as a key benefit of interpretive linking. P5 noted that small discrepancies in content can be damaging: *“there’s nothing worse than when a number on a slide doesn’t match the number in the chart.”* P5 also reflected that even *“a small change, like modifying a single number, might have broader analytical implications.”* and explained that AI could help surface these deeper implications. At the same time, participants noted that propagating changes automatically introduces new burdens. P2 reflected that tracking changes with cascading effects proved tedious, and P5 thought they might be *“jarring for other users”* in collaborative workflows.

Concerns about failure amplified these tradeoffs. P3 emphasized that even a low rate of error in propagated changes could have an outsized impact on trust and adoption, noting that *“even a small failure rate would massively affect my confidence in the system.”* Together, these reflections frame change propagation as a design problem—one that requires careful attention to visibility, trust, and user control alongside the efficiencies interpretive linking may offer.

A different interaction paradigm with AI. Participants highlighted how interpretive linking suggests a different interaction paradigm for generative AI, distinct from chat-based prompting. P4 and P6 noted a reduced need for explicit prompts, with P4 observing that *“prompting is huge [difficult] — people don’t know what to write.”* Instead, interaction shifted toward specifying desired outcomes and directly editing AI-generated content. P4 described this as a *“break in the existing mental model of how to communicate with*

LLMs,” emphasizing the value of being able to *“edit right away and focus on the content, not on how to talk to the system.”* Together, these accounts position interpretive linking as a mediator for coordinating meaning across representations, supporting reflection, collaboration, and iterative sensemaking with generative AI.

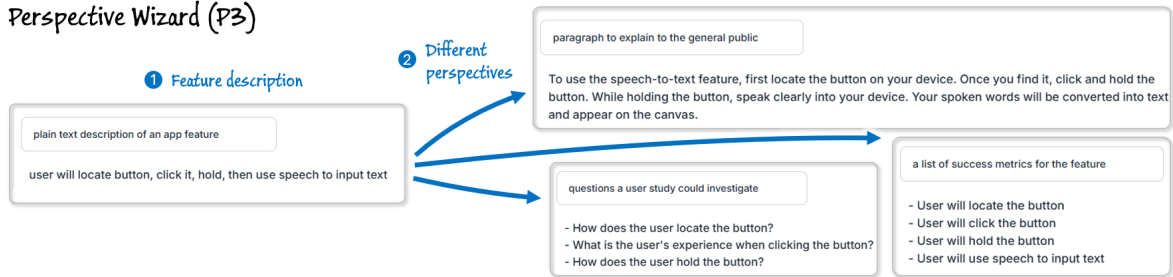
6 Discussion and Future Work

We proposed *interpretive linking* as a general approach that extends brushing and linking to knowledge work. Unlike traditional linking based on stable identifiers and deterministic mappings, interpretive linking coordinates evolving, ambiguous relationships between representations using generative AI to infer and translate change, with the goal of preserving both semantic coherence and the communicative intent of individual artifacts. Our technology probe demonstrates that such AI-mediated linking is feasible and our study suggests that it may make repurposing material more fluid, enabling authors to work with different representations without manually reconstructing relationships or propagating edits.

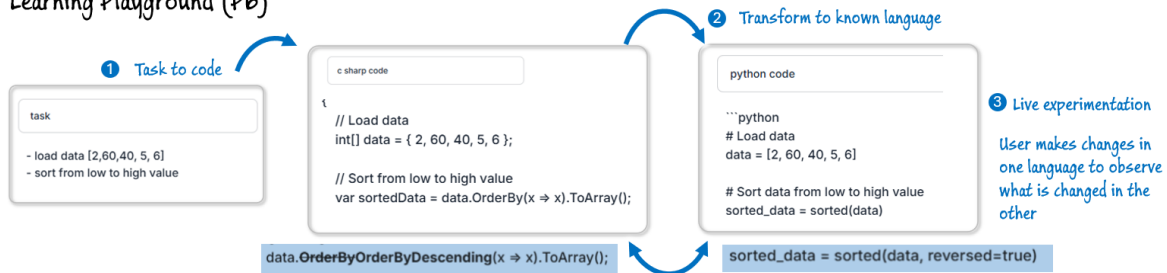
New Opportunities. Participants articulated workflow ideas in which interpretive linking could support *reflection and ideation* through alternative representations. By working on the same content in different forms, participants explored how ideas co-evolved across representations. These workflows allowed participants to reason about their work by moving fluidly between different lenses as part of the thinking and ideation process. Participants also described opportunities for interpretive linking to help *coordinate work across collaborators* with different expertise and communicative practices. By linking representations tailored to each collaborator (e.g., accessibility needs, domain expertise, form of expression), interpretive linking supported collaboration beyond having a single shared artifact. More broadly, reflections from participants suggest interpretive linking as *an alternative way of engaging with generative AI*. Rather than treating AI primarily as a separate tool for producing outputs, participants described interpretive linking as a means for interacting with AI by directly editing content and meta-descriptions rather than formulating instructions and copy-pasting content in a separate chat application.

Open Challenges. At the same time, our probe revealed significant challenges, particularly around propagating changes across modalities. While change propagation was effective within a single modality, text-to-image propagation exposed tensions between control and flexibility. Image generation models often introduced subtle visual changes or reinterpreted content based on transformation history and prompt details. More controlled pipelines using image

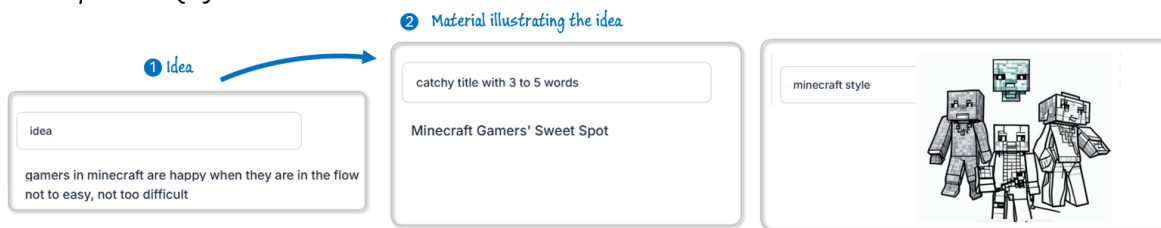
Perspective Wizard (P3)



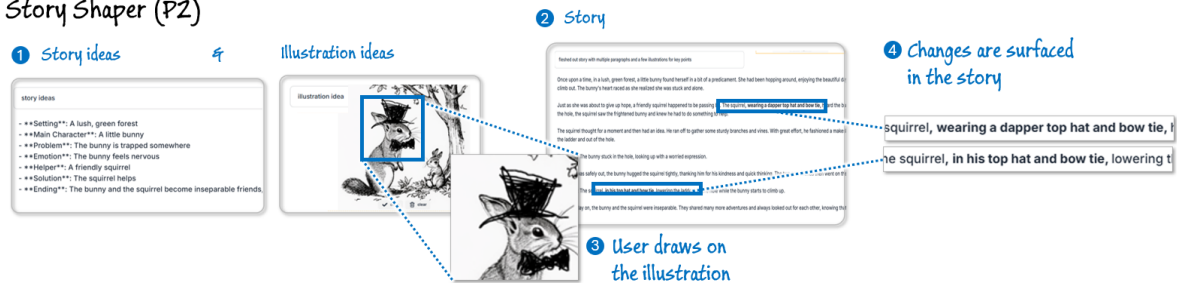
Learning Playground (P6)



Idea Crystallizer (P1)



Story Shaper (P2)



Living TLDR (P5)

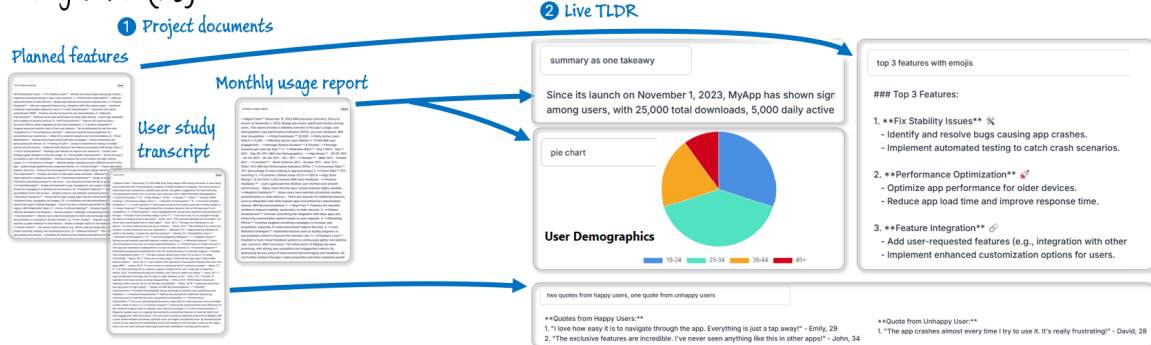


Figure 5: Probe-to-prototype annotated screen captures for five of our six participants.

segmentation and inpainting increased predictability but raised new design questions, such as how to relax constraints or adapt masks (e.g., whether a mask should be expanded when a replaced object is larger than the original, as in the elephant-for-bunny example in Figure 4). As pipelines became more prescriptive, they eroded a key strength observed in text-based transformations: the ability to operate with minimal specification and generalize across diverse contexts. These findings suggest that no single pipeline suffices for interpretive linking today. Effective propagation requires strategies tuned to specific scenarios, modalities, and representational goals, often balancing inference and constraint. More generally, potential semantic drift that accumulates over frequent regeneration of linked content remains a challenge: when changes are propagated across chained representations, meta-description constraints can weaken from iteration to iteration, gradually pulling the target away from the source's communicative intent. In high-stakes scenarios, achieving reliable and accurate AI-mediated transformations of content as part of interpretive linking is critical, which remains a challenge for current state-of-the-art generative models.

Future Work. While our qualitative study involved a relatively small number of participants, it surfaced a rich set of workflow ideas and scenarios that illuminate how interpretive linking might be appropriated in practice. But the sample size of six participants and their similar background (working in technology and design roles within a single large organization), while appropriate for an exploratory probe study, limit generalizability. Studying interpretive linking with knowledge workers in other domains (e.g., law, science, finance, and education) is an important direction for future work, alongside larger and more longitudinal studies that examine how interpretive linking is appropriated over time.

Participants' reflections also highlighted key challenges, including managing the scope, cost, and risk of propagating changes, ensuring reliability and traceability, and balancing automation with user control. Addressing these issues may require interaction mechanisms that make change visible and negotiable, such as support for inspection, versioning, and affording reversibility, as well as more explicit ways of specifying scope of propagation and negotiating ambiguity across representations. More broadly, our findings point to open questions about whether interpretive linking is best supported through pipelines tightly integrated into specific workflows or through more generalizable approaches that span domains and modalities. Technical challenges, including model latency, computational cost, and support for additional modalities or finer-grained transformations, also remain areas for future work.

References

- [1] Richard A Becker and William S Cleveland. 1987. Brushing scatterplots. *Technometrics* 29, 2 (1987), 127–142. <https://doi.org/10.1080/00401706.1987.10488204>
- [2] Virginia Braun and Victoria Clarke. 2022. *Thematic Analysis: A Practical Guide*. SAGE, Los Angeles. <https://login.eux.idm.oclc.org/login?url=https://resolver.vitalsource.com/9781526417299>
- [3] Frederik Brudy, Christian Holz, Roman Rädle, Chi-Jui Wu, Steven Houben, Clemens Nylandsted Klokmoose, and Nicolai Marquardt. 2019. Cross-Device Taxonomy: Survey, Opportunities and Challenges of Interactions Spanning Across Multiple Devices. In *Proceedings of the 2019 CHI Conference on Human Factors in Computing Systems*. ACM, Glasgow Scotland UK, 1–28. <https://doi.org/10.1145/3290605.3300792>
- [4] Yining Cao, Yiyi Huang, Anh Truong, Hujung Valentina Shin, and Haijun Xia. 2025. Compositional Structures as Substrates for Human-AI Co-creation Environment: A Design Approach and A Case Study. In *Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems (CHI '25)*. Association for Computing Machinery, New York, NY, USA, Article 188, 25 pages. <https://doi.org/10.1145/3706598.3713401>
- [5] Stuart K Card, Jock Mackinlay, and Ben Shneiderman. 1999. *Readings in information visualization: using vision to think*. Morgan Kaufmann.
- [6] Bay-Wei Chang, Jock D. Mackinlay, Polle T. Zellweger, and Takeo Igarashi. 1998. A negotiation architecture for fluid documents. In *Proc. of UIST '98* (San Francisco, California, USA). ACM, New York, NY, USA, 123–132. <https://doi.org/10.1145/288392.288585>
- [7] Fanny Chevalier, Pierre Dragicevic, and Christophe Hurter. 2012. Histomages: fully synchronized views for image editing. In *Proc. of UIST '12* (Cambridge, Massachusetts, USA). ACM, New York, NY, USA, 281–286. <https://doi.org/10.1145/2380116.2380152>
- [8] James M Clark and Allan Paivio. 1991. *Dual coding theory and education*. Educational Psychology Review. <https://doi.org/10.1007/BF01320076>
- [9] Niklas Elmqvist, Pierre Dragicevic, and Jean-Daniel Fekete. 2008. Rolling the Dice: Multidimensional Visual Exploration using Scatterplot Matrix Navigation. *IEEE Transactions on Visualization and Computer Graphics* 14, 6 (2008), 1141–1148. <https://doi.org/10.1109/TVCG.2008.153>
- [10] Niklas Elmqvist, Nathalie Henry, Yann Riche, and Jean-Daniel Fekete. 2008. Melange: space folding for multi-focus interaction. In *Proc. of CHI '08* (Florence, Italy). ACM, New York, NY, USA, 1333–1342. <https://doi.org/10.1145/1357054.1357263>
- [11] Jessica L Feuston and Siân E Lindley. 2018. How social dynamics and the context of digital content impact workplace remix. In *Proc. of CHI '18*. 1–13. <https://doi.org/10.1145/3173574.3174171>
- [12] Deepanway Ghosal, Navonil Majumder, Ambuj Mehrish, and Soujanya Poria. 2023. Text-to-Audio Generation using Instruction-Tuned LLM and Latent Diffusion Model. arXiv:2304.13731 [eess.AS] <https://arxiv.org/abs/2304.13731>
- [13] Ken Hinckley, Gonzalo Ramos, Francois Guimbretiere, Patrick Baudisch, and Marc Smith. 2004. Stitching: pen gestures that span multiple displays. In *Proceedings of the working conference on Advanced visual interfaces*. ACM, Gallipoli Italy, 23–31. <https://doi.org/10.1145/989863.989866>
- [14] Steven Houben and Nicolai Marquardt. 2015. WatchConnect: A Toolkit for Prototyping Smartwatch-Centric Cross-Device Applications. In *Proceedings of the 33rd Annual ACM Conference on Human Factors in Computing Systems*. ACM, Seoul Republic of Korea, 1247–1256. <https://doi.org/10.1145/2702123.2702215>
- [15] Hilary Hutchinson, Wendy Mackay, Bo Westerlund, Benjamin B. Bederson, Allison Druiin, Catherine Plaisant, Michel Beaudouin-Lafon, Stéphane Conversy, Helen Evans, Heiko Hansen, Nicolas Roussel, and Björn Eiderbäck. 2003. Technology probes: inspiring design for and with families. In *Proc. of CHI '03* (Ft. Lauderdale, Florida, USA). ACM, New York, NY, USA, 17–24. <https://doi.org/10.1145/642611.642616>
- [16] Carlos Jensen, Heather Lonsdale, Eleanor Wynn, Jill Cao, Michael Slater, and Thomas G Dietterich. 2010. The life and times of files and information: a study of desktop provenance. In *Proc. of CHI '10*. 767–776. <https://doi.org/10.1145/1753326.1753439>
- [17] Albert Q. Jiang, Alexandre Sablayrolles, Arthur Mensch, Chris Bamford, Devendra Singh Chaplot, Diego de las Casas, Florian Bressand, Gianna Lengyel, Guillaume Lample, Lucile Saulnier, Léo Renard Lavaud, Marie-Anne Lachaux, Pierre Stock, Teven Le Scao, Thibaut Lavril, Thomas Wang, Timothée Lacroix, and William El Sayed. 2023. Mistral 7B. arXiv:2310.06825 [cs.CL] <https://arxiv.org/abs/2310.06825>
- [18] Clemens N. Klokmoose, James R. Eagan, Siemen Baader, Wendy Mackay, and Michel Beaudouin-Lafon. 2015. *Webstrates: Shareable Dynamic Media*. In *Proceedings of the 28th Annual ACM Symposium on User Interface Software & Technology*. ACM, Charlotte NC USA, 280–290. <https://doi.org/10.1145/2807442.2807446>
- [19] Siân Lindley, Chris Elsdén, Hanna Moser, and Amid Ayobi. 2023. Content Repurposing in Knowledge Work: Implications for Generative AI. In *Presented at CHI 2023 Workshop on Generative AI and HCI*.
- [20] Richard E. Mayer. 2009. *Multimedia Learning* (2 ed.). Cambridge University Press. <https://doi.org/10.1017/CBO9780511811678>
- [21] Chris North and Ben Shneiderman. 2000. Snap-together visualization: a user interface for coordinating visualizations via relational schemata. In *Proceedings of the working conference on Advanced visual interfaces*. 128–135. <https://doi.org/10.1145/345513.345282>
- [22] Xiaohan Peng, Janin Koch, and Wendy E. Mackay. 2024. DesignPrompt: Using Multimodal Interaction for Design Exploration with Generative AI. In *Proceedings of the 2024 ACM Designing Interactive Systems Conference (Copenhagen, Denmark) (DIS '24)*. Association for Computing Machinery, New York, NY, USA, 804–818. <https://doi.org/10.1145/3643834.3661588>
- [23] Jonathan C Roberts. 2007. State of the art: Coordinated & multiple views in exploratory visualization. In *Fifth international conference on coordinated and multiple views in exploratory visualization (CMV 2007)*. IEEE, 61–71. <https://doi.org/10.1109/CMV.2007.20>
- [24] Robin Rombach, Andreas Blattmann, Dominik Lorenz, Patrick Esser, and Björn Ommer. 2021. High-Resolution Image Synthesis with Latent Diffusion Models. arXiv:2112.10752 [cs.CV]

- [25] Mark Sadoski and Allan Paivio. 2013. *Imagery and Text: A Dual Coding Theory of Reading and Writing*. Routledge. <https://doi.org/10.4324/9780203801932>
- [26] John Stasko. 2014. Value-driven evaluation of visualizations. In *Proceedings of the Fifth Workshop on Beyond Time and Errors: Novel Evaluation Methods for Visualization*. 46–53. <https://doi.org/10.1145/2669557.2669579>
- [27] Jeffrey Stylos, Brad A Myers, and Andrew Faulring. 2004. Citrine: providing intelligent copy-and-paste. In *Proc. of UIST '04*. 185–188. <https://doi.org/10.1145/1029632.1029665>
- [28] Sangho Suh, Meng Chen, Bryan Min, Toby Jia-Jun Li, and Haijun Xia. 2024. Luminate: Structured Generation and Exploration of Design Space with Large Language Models for Human-AI Co-Creation. In *Proceedings of the CHI Conference on Human Factors in Computing Systems*. 1–26. <https://doi.org/10.1145/3613904.3642400> arXiv:2310.12953 [cs].
- [29] John Sweller. 1988. Cognitive Load During Problem Solving: Effects on Learning. *Cognitive Science* 12, 2 (1988), 257–285. https://doi.org/10.1207/s15516709cog1202_4
- [30] Peter Tandler, Thorsten Prante, Christian Müller-Tomfelde, Norbert Streitz, and Ralf Steinmetz. 2001. Connectables: dynamic coupling of displays for the flexible creation of shared workspaces. In *Proceedings of the 14th annual ACM symposium on User interface software and technology*. ACM, Orlando Florida, 11–20. <https://doi.org/10.1145/502348.502351>
- [31] Ye Tian, Ling Yang, Haotian Yang, Yuan Gao, Yufan Deng, Jingmin Chen, Xintao Wang, Zhaochen Yu, Xin Tao, Pengfei Wan, Di Zhang, and Bin Cui. 2024. VideoTetris: Towards Compositional Text-to-Video Generation. arXiv:2406.04277 [cs.CV] <https://arxiv.org/abs/2406.04277>
- [32] Hugo Touvron, Thibaut Lavril, Gautier Izacard, Xavier Martinet, Marie-Anne Lachaux, Timothée Lacroix, Baptiste Rozière, Naman Goyal, Eric Hambro, Faisal Azhar, Aurelien Rodriguez, Armand Joulin, Edouard Grave, and Guillaume Lample. 2023. LLaMA: Open and Efficient Foundation Language Models. arXiv:2302.13971 [cs.CL] <https://arxiv.org/abs/2302.13971>
- [33] Priyan Vaithilingam, Elena L. Glassman, Jeevana Priya Inala, and Chenglong Wang. 2024. DynaVis: Dynamically Synthesized UI Widgets for Visualization Editing. <http://arxiv.org/abs/2401.10880> arXiv:2401.10880 [cs].
- [34] Sitong Wang, Samia Menon, Tao Long, Keren Henderson, Dingzeyu Li, Kevin Crowston, Mark Hansen, Jeffrey V Nickerson, and Lydia B Chilton. 2024. Reel-Framer: Human-AI Co-Creation for News-to-Video Translation. In *Proceedings of the 2024 CHI Conference on Human Factors in Computing Systems* (Honolulu, HI, USA) (CHI '24). Association for Computing Machinery, New York, NY, USA, Article 169, 20 pages. <https://doi.org/10.1145/3613904.3642868>
- [35] Zheng Zhang, Jie Gao, Ranjodh Singh Dhaliwal, and Toby Jia-Jun Li. 2023. VISAR: A Human-AI Argumentative Writing Assistant with Visual Programming and Rapid Draft Prototyping. In *Proceedings of the 36th Annual ACM Symposium on User Interface Software and Technology*. 1–30. <https://doi.org/10.1145/3586183.3606800> arXiv:2304.07810 [cs].
- [36] Chen Zhu-Tian, Zeyu Xiong, Xiaoshuo Yao, and Elena Glassman. 2024. Sketch Then Generate: Providing Incremental User Feedback and Guiding LLM Code Generation through Language-Oriented Code Sketches. <https://doi.org/10.48550/arXiv.2405.03998> arXiv:2405.03998 [cs].
- [37] John Zimmerman, Jodi Forlizzi, and Shelley Evenson. 2007. Research through design as a method for interaction design research in HCI. In *Proceedings of the SIGCHI conference on Human factors in computing systems*. 493–502. <https://doi.org/10.1145/1240624.1240704>