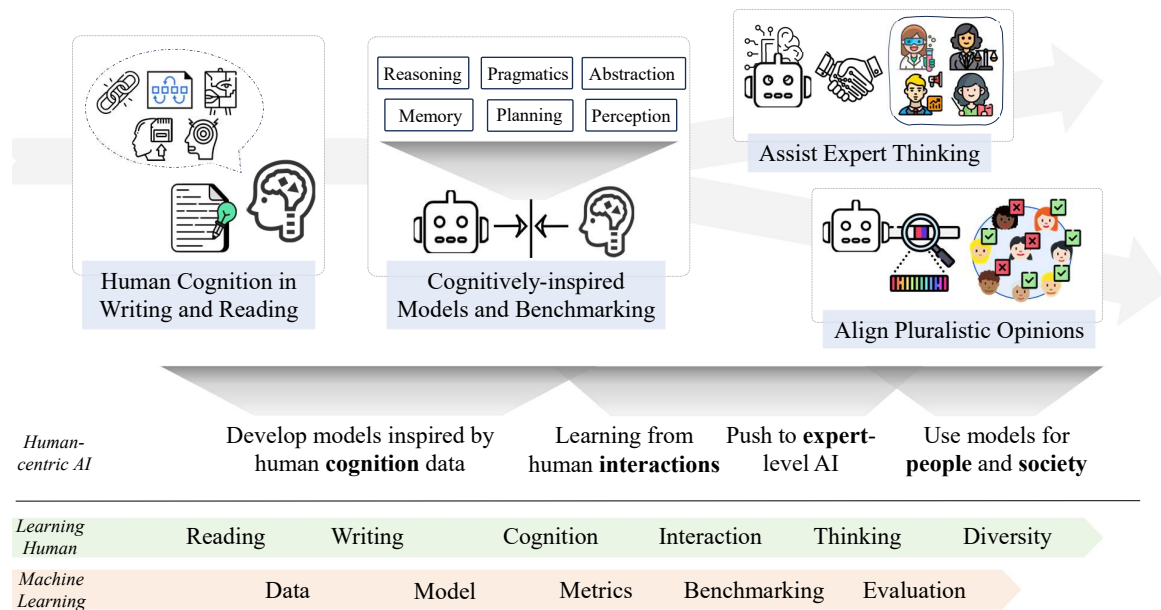


AI as Thinking Partners: Supporting Experts' Cognitive Workflows

Dongyeop Kang

I aim to build human-centric language technologies that cognitively align human and machine thinking, advancing AI as thinking partners for experts. Today's language models still struggle with long-horizon, iterative, self-reflective reasoning, which is essential for high-fidelity decision making in real-world professional work. My work bridges language and cognition to develop novel algorithms, benchmarks, and interaction frameworks that support experts in complex, real-world cognitive workflows. What differentiates my agenda is its training signal: models learn not only from the final text people produce, but from the cognitive processes behind it—keystroke-level revisions, gaze, and expert workflow traces. I design AI that augments humans: models that learn from how people plan, reason, and create; anticipate cognitive bottlenecks; scaffold difficult tasks; and adapt dynamically to expert strategies.



The central premise of my work is that AI should extend and enhance human cognition: supporting expert reasoning, reducing cognitive burdens, and accommodating diverse individual goals. I organize my research into three tightly interconnected pillars: (1) *Cognitive Scaffolding*—collecting human cognition data (§1), building cognitively-aligned AI models (§2), and benchmarking human and machine cognition (§3); (2) *Thinking Assistants for Experts* (§4); and (3) *Societal Alignment* through pluralistic AI (§5).

§1 Collecting Human Cognition Data: Human cognition involves complex, multi-layered thinking processes. For example, writing is a non-linear and iterative cognitive process. By analyzing human writing-process data—such as our *ScholaWrite* corpus of keystroke-level scholarly writing [71] and iterative revision datasets [22, 62, 79]—we can gain insights into these thought processes, which can then be used to improve the planning and reasoning capabilities of AI models [90, 21]. Similarly, human perception can be studied through reading behaviors, such as eye-tracking data or explicit perception annotations, offering clues about how people perceive and process information through reading [30, 14, 34]. Understanding both writing and reading behaviors enables us to enhance the cognitive capabilities of LLMs and develop AI assistants that better support human thinking.

§2 Cognitively-aligned AI Models: Cognitive alignment seeks to extend LLM capabilities to complex human behaviors, including neurosymbolic reasoning [52], task compositionality [90, 29], hierarchical planning [43, 48], abstraction, and social cognition [15]. We develop novel learning paradigms that either mimic human cognitive processes directly or incorporate cognitive data (e.g., writing, reading, expert reasoning, interaction) into training objectives. This includes self-supervised planning, structural and multi-attribute alignment, self-evolving reward models (e.g., Meta Policy Optimization

[60]), and cognitively-efficient modeling strategies.

§3 Cognitive Benchmarking: Current LLM evaluation largely relies on shallow, monotonic tasks. We address this through three directions: (i) assessing both the potential [78, 100, 56, 57, 88, 55, 3] and risks [67, 11, 62, 79] of AI-generated data and AI-based evaluation, identifying issues such as cognitive bias, stylistic artifacts, and shallow synthesis; (ii) building *CogBench*, a large-scale cognitive assessment framework [18, 17] that integrates rigorous cognitive science methodologies with scalable evaluation to contrast human and machine cognition, in open collaboration with interdisciplinary researchers; and (iii) developing long-horizon behavioral evaluation for agentic systems [61, 82].

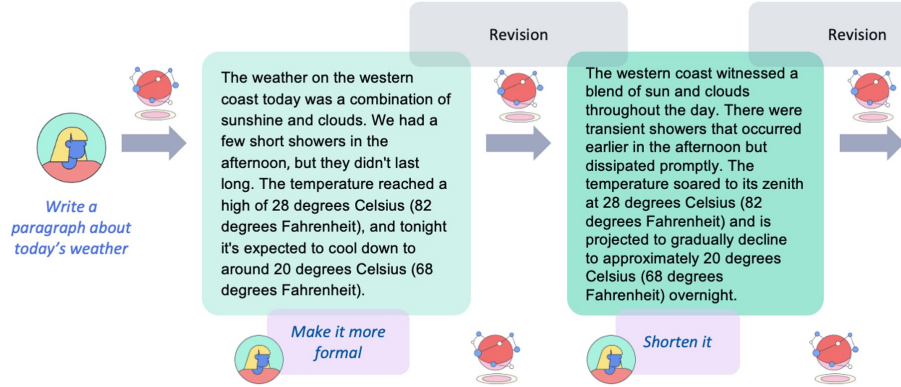
§4 Thinking Assistants for Experts: In domains such as science and law, there is a significant cognitive gap between human experts and current LLMs. We design interactive, domain-specific AI systems that facilitate productive collaboration, adapt to expert workflows, and provide cognitively-aligned assistance. For scientists, we have developed reading tools like *Semantic Reader* [34, 77] and *SciTalk* [87], as well as writing assistants that detect overloaded symbols and discourse-level inconsistencies and promote iterative refinement [21, 71]. For legal professionals, we built *LawFlow* [12], a dataset capturing lawyers’ long-horizon planning and reasoning, and are building legal assistants on top of it.

§5 Societal Alignment via Pluralism: AI technologies must reflect diverse human perspectives to achieve societal alignment. We develop data-centric methods to detect, characterize, and augment underrepresented viewpoints [78, 95, 57, 54], and model-centric methods to encode pluralism at individual, group, and societal levels using distributional alignment and societal value modeling [55, 31, 15]. This mitigates risks of monolithic outputs and supports socially inclusive AI.

Together, these agendas form my long-term vision: *human-centric thinking partners* that perform complex, multi-step reasoning, reduce cognitive load, and reflect the diversity and values of human society. My research integrates linguistics, social sciences, and cognitive sciences, and is supported by industry and government partners such as Grammarly, Naver, Google, NSF, Cisco, Sony, Accenture, Thomson Reuters, and Open Philanthropy, with active research collaborations with Amazon, AI2, Naver, and Google. I hold affiliations with cross-college labs at UMN and collaborate widely with faculty and students in computer science, law, psychology, education, journalism, design, and medicine. I also foster cross-community collaboration through workshop organization: I co-organized the first “CtrlGen: Controllable Generative Modeling” workshop at NeurIPS 2021, and founded the “In2Writing: Intelligent and Interactive Writing Assistants” workshop series at ACL 2022, CHI 2023, and CHI 2024, connecting the ML, HCI, NLP, and professional writing communities. I also founded the “Pluralistic Alignment” workshop at NeurIPS 2024 to emphasize diversity in AI.

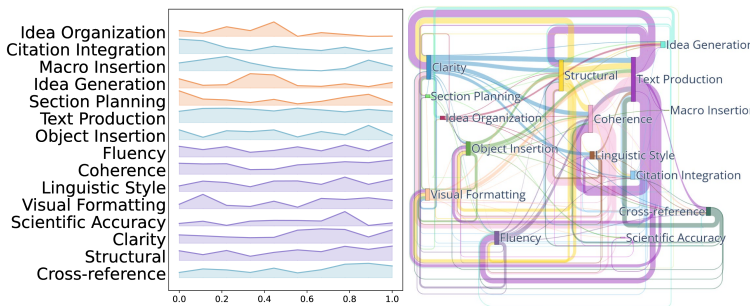
1 Collecting Human Cognition Data

Language serves as a window into human cognition. Writing and reading are not linear acts but complex, iterative processes that reveal how people plan, reason, and perceive information. For example, a well-written manuscript reflects a writer’s long-term, multi-stage thinking, far beyond simple next-token prediction. Similarly, reading behaviors captured through eye-tracking or annotated perceptions provide insights into how individuals interpret, learn from, and engage with text. This first research direction—the data foundation of my *Cognitive Scaffolding* pillar—focuses on systematically collecting and analyzing human cognition data from both *production* (writing) and *perception* (reading) to better understand these processes and inform cognitively-aligned AI systems. We call this **cognitive scaffolding**: the systematic use of temporally grounded human cognitive traces—rather than isolated final outputs—as training signals.

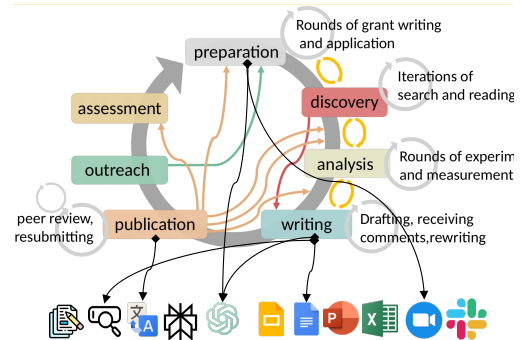


1.1 Writing Data

Human writing involves richer and more deliberate processes than next-token generation: drafting, revising, planning, proofreading, and synthesizing knowledge—often collaboratively. Domain-specific writing, such as drafting scientific manuscripts, adds further challenges, requiring sustained practice, peer interaction, and integration of prior literature. We first examine the *iterative nature of human text revision*. We have collected revision data across domains (*IteraTeR* [22, 59]) and built systems that emulate iterative refinement until quality stabilizes (*R3* [21], *CoEdIT* [90]). Human–AI collaborative editing consistently outperforms either humans or AI working alone [21, 90, 73]. Our studies show that explicitly modeling the revision process—whether through AI alone or human–AI collaboration—consistently yields higher text quality, demonstrating the effectiveness of test-time scaling. Guiding models with human revision intents (e.g., clarity, coherence) further produces revision traces that better mimic the human writing process, making these models better suited to supporting human writers. We also study how varying levels of LLM access reshape writers’ behavior [8].



(a) Change of writing intents in scholarly writing [71]



(b) Research Workflow with AI Tools

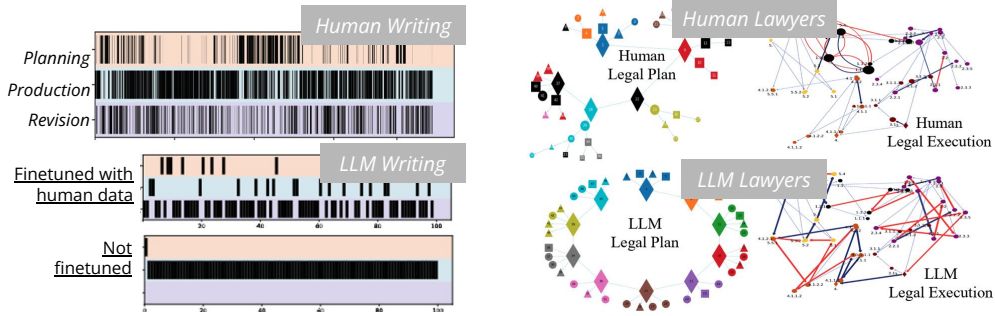
Instead of optimizing isolated skills, we model the coordinated dynamics of full expert processes. In *ScholaWrite* [68, 71], we captured end-to-end scientific writing trajectories—nearly 62K keystroke-level text changes with intention annotations, collected over four months of real manuscript writing—and showed that training on human process data improves alignment and writing quality. This dataset

provides a strong empirical basis for cognitively-aligned writing assistants that model the full lifecycle of text production.

Building on this, the *ScholaWrite2* project is collecting year-long longitudinal research workflows from scientists across domains, including keystrokes, screen and room video from XR glasses, and meeting notes. The collected workflows will be annotated with task intent, cognitive state, and AI-interaction signals, with the goal of formalizing an ontology of real-world scientific reasoning. The resulting corpus forms an **ontology of the scientific process**, enabling training of agents and interfaces that reflect real workflows; relatedly, we find that structured constraints on sensemaking yield more novel research output [80]. Ultimately, we aim to establish a longitudinal repository of cognitive workflows (science, law, education), allowing models to learn not isolated tasks but the temporal dynamics of human reasoning.

1.2 Long-horizon Workflow Data from Domain Experts

My long-term goal is cognitively inspired AI that performs complex, multi-step tasks involving causality, abstraction, and memory while reducing human cognitive load. Such systems require multiple cognitive functions working in concert across an entire workflow (e.g., scientific processes, legal reasoning), not excellence in any single capability. I therefore collect longitudinal workflows from both human experts and LLMs, and use them to improve models' cognitive capabilities on complex, high-fidelity tasks in knowledge-intensive domains.



Left: Scholarly Writing Intents over Time [71], Right: Legal Planning (left) & Execution (right) [12]

Comparing Human and LLM Workflows. The first step is to investigate how human workflows differ from AI workflows and to design methods to align them. Our prior work on collecting long-term human workflow data (in scholarly writing [71] and legal processing [12]) has revealed clear distinctions between human and AI processes. When end-to-end writing data is used to train models to emulate human writing trajectories, the resulting systems exhibit more human-aligned dynamics and improved writing quality [71, 68]. Similarly, in legal planning and execution, we find that humans often act more spontaneously, sometimes repeating inefficient steps [12], whereas AI tends to rigidly follow pre-defined plans. We observe these contrasting behavioral patterns in our study of human- and AI-based legal workflows with practicing lawyers [92], providing key insights for designing cognitively-aligned thinking assistants that better bridge workflows between lawyers and AI agents.

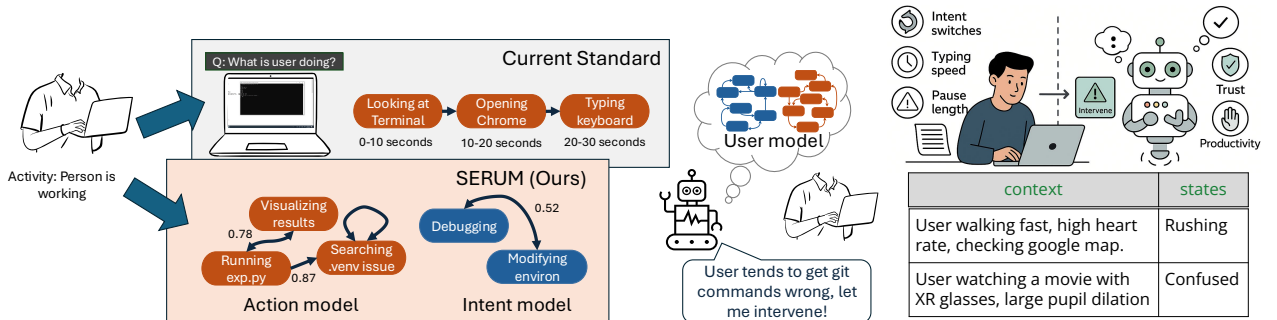
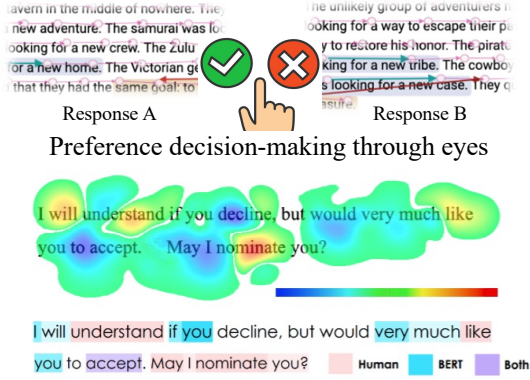


Figure. 2: (Left) SERUM's [89] multi-pass pipeline revisits prior context across frames, enabling the construction of a refined user model for both user actions and intents, enabling better informed proactive suggestions. (Right) Psychometrics from user workflow

User Cognitive Behaviors from Multimodal Context. To scale workflow acquisition beyond keystroke logging and annotation, we are developing multimodal workflow capture and cognitive state modeling. We built systems that extract fine-grained workflow dynamics from raw egocentric YouTube videos (e.g., programming, cooking, tutorials, market research) and infer latent intents using theory-grounded psychometrics. Our platform induces probabilistic finite user state machines (FUSMs)—compact graphs of user actions and the latent intents driving transitions between them—from workflow video via iterative, bottom-up action and intent induction, and derives cognitive metrics from multimodal signals such as gaze, physiology, voice, and video from XR devices. We plan to train lightweight cognitive reasoning models to infer hidden states online and adapt assistance in real time; our *SERUM* framework takes a first step by extracting and refining user states from interaction data [89].

1.3 Reading Data

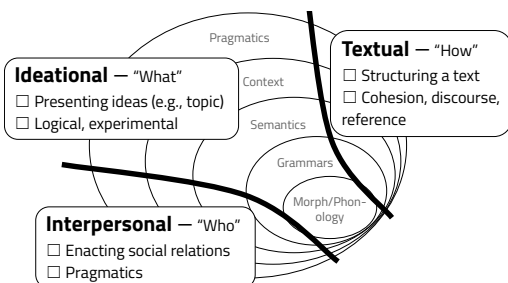


Understanding *politeness* through eyes or annotations

Metric	Illustration	Interpretation
Prompt reread		Indicates uncertainty or careful prompt reconsideration
Response loop		Strong indecision between candidate responses
Read time		Higher values indicate deeper response evaluation
Word coverage		Lower coverage indicates careless reading or a clear-cut decision

Whereas writing data reveals the generative processes of human thought, reading data captures the perceptual and evaluative side of cognition—how people interpret, prioritize, and react to information. We have collected explicit annotations of stylistic understanding [30] and used them to train LLMs that align with readers’ stylistic judgments [32]. We have also gathered eye-tracking data to study how readers engage with stylistic cues, narrative content, or preference-ranking annotation tasks [14, 85, 19], showing that gaze patterns provide richer cognitive signals than explicit annotations or model-based interpretations. As shown in the table above, these fine-grained gaze patterns also provide interpretable signals of user behaviors. By collecting data that closely mirrors human perception and evaluative behavior, we aim to train AI systems that more faithfully reproduce human-like decision making and perception processes [55]. These results show that perceptual signals capture evaluative cognition that is systematically absent from text-only supervision; a scaffolding framework for recovering such lost context is developed in [9].

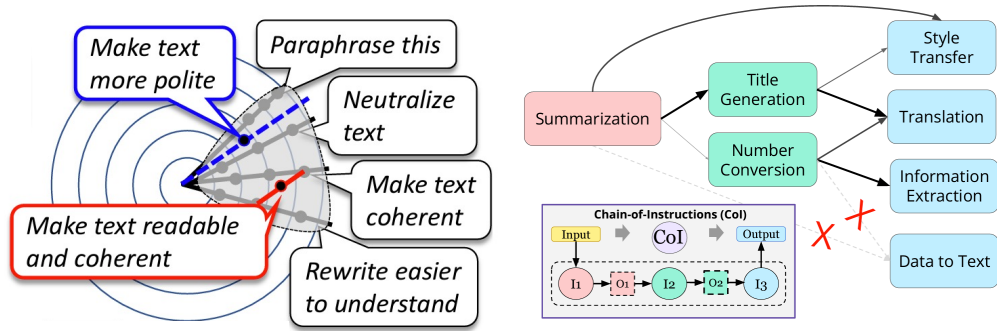
2 Cognitively-aligned AI Models



The goal of *cognitive alignment* is to enhance LLMs’ ability to emulate complex human behaviors—such as task compositionality, planning, abstraction, reasoning, and memory—by either (i) mimicking cognitive processes directly or (ii) incorporating human cognition data into training objectives. In line with Halliday’s Systemic Functional Linguistics (SFL)—a theory that language simultaneously encodes content, social relations, and structure [41]—these efforts address a single question: how can models internalize human-like structure across reasoning, planning, social judgment, and efficiency?

2.1 Reasoning and Compositionality

Human reasoning often bridges incomplete information through external knowledge and iterative refinement. We integrate neural and symbolic systems [52, 23, 51] to fill knowledge gaps while enhancing interpretability [45]. To test and improve compositionality, we created benchmarks and models for composite and chained tasks—including *CoEdit*, our publicly released instruction-tuned

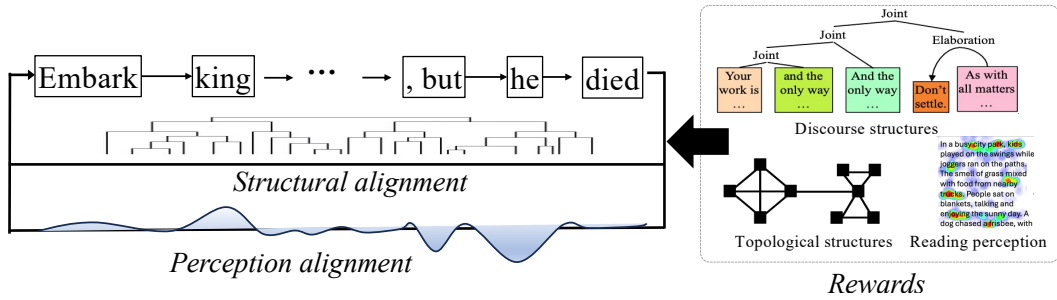


text-editing system [90], and *Chain-of-Instructions* [29]—using compositional data augmentation and task-space densification, i.e., synthesizing training instances for unseen combinations of atomic tasks [6]. Human reasoning is often iterative, refining ideas or solutions step-by-step. We mimic this by developing models that emulate iterative reasoning patterns, showing that step-by-step improvements lead to higher-quality, more diverse outputs [22, 59, 21, 36].

2.2 Planning

Humans achieve coherence by organizing content into a unified structure through decisions about topics, ordering, abstraction, and communication strategy. We therefore model planning as hierarchical optimization of text generation, coordinating multiple text units to guide surface realization.

Discourse-Guided Planning. Coherence emerges from plans that align sentence sequences with discourse goals, grounded in rhetorical structure theory and script theory. We develop supervised planners that incorporate discourse relations [45, 46], topical keywords [48], and social or persuasive goals [28]. Explicit plans based on relations, keywords, and goals enable more controlled, interpretable, and coherent generation through theory-driven discourse guidance [64, 76, 28].



Self-Supervised Planning and Structural Alignment. Large language models rely on next-token prediction, which creates a gap in complex planning and long-horizon writing. We address this gap using alignment methods such as reinforcement learning to optimize high-level policies with human or model-based feedback. Our self-supervised role-playing framework [43] applies reinforcement learning to simulate dialogue interactions and optimize explicit communication goals. This paradigm has evolved into reinforcement learning with AI feedback, where one agent generates text and another provides structured feedback, supporting long-term policy learning. Recent work further aligns text structure with human writing through structural rewards [65, 62, 93, 102] and integrates human perception signals via dense reward alignment, assigning rewards at the span level rather than to whole responses [69].

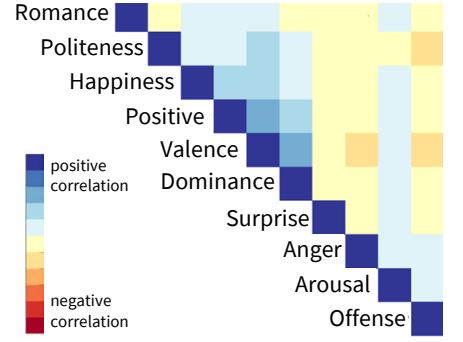
2.3 Social Cognition

Language encodes interpersonal and societal dimensions of social cognition. Text style emerges from interacting factors such as formality, emotion, and metaphor, reflecting an author’s personality and serving specific communicative goals. Understanding these cross-style relationships is key to capturing the nuances of human communication. We built datasets for cross-style analysis, including PASTEL [44] and xSLUE [49], showing that multi-style learning outperforms single-style approaches. Certain style pairs, such as impoliteness–offense, are strongly correlated, while contradictory combinations yield less appropriate outcomes. This underexplored area remains vital

for modeling complex stylistic patterns. More recently, we expanded to high-level affective states such as nostalgia [58] and skepticism [83], collecting annotated data for social cognition research. To align models with multiple social aspects, we developed methods for multi-attribute control via data balancing [13], multi-task fine-tuning [49], policy learning with dynamic reward re-weighting [15], or distillation [7].

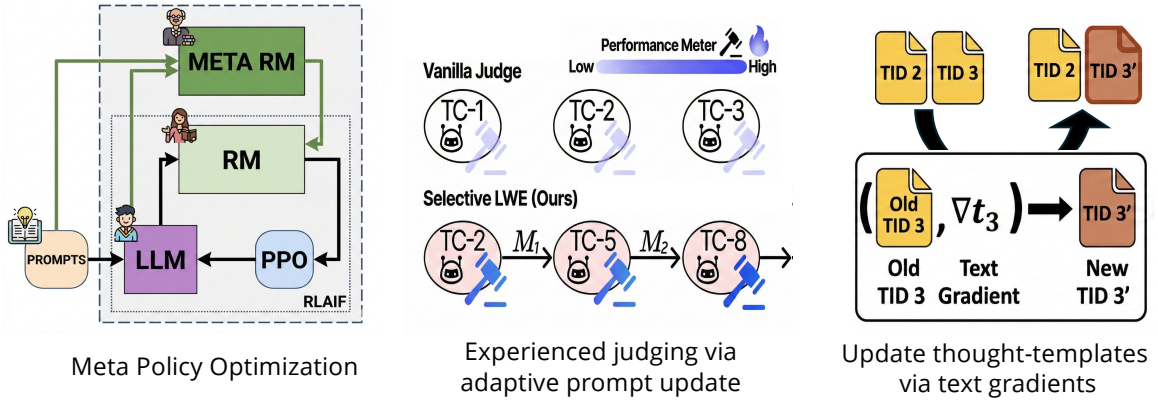
2.4 Cognitive Efficiency

Abstraction enables humans to process knowledge efficiently, bypassing detailed reasoning when shortcuts suffice. Inspired by this, we design algorithms that optimize data usage, computation, and memory in LLM training, reducing hallucinations and improving reliability. Our work includes data-efficient methods that maintain performance with fewer training samples [57, 88], and compute-efficient models that pair parameters across models via dynamic weight warping [81]. We continue to develop cognitively inspired techniques that optimize LLMs’ *cognitive utility function*, balancing performance with resource efficiency [96].



2.5 Self-Evolving Agents with Adaptive Cognitive Scaffolding

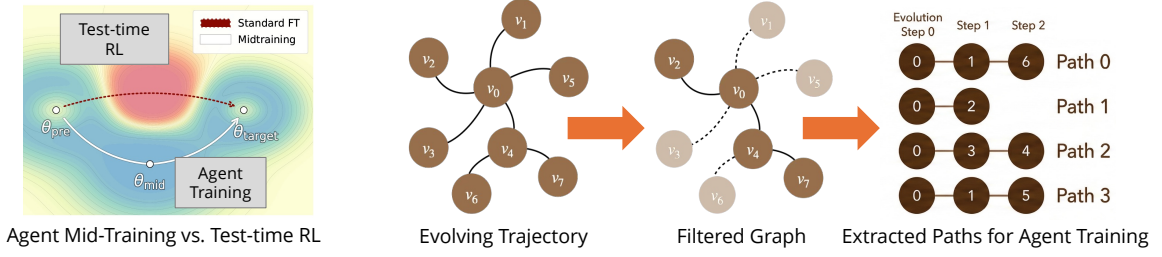
Beyond modeling human workflows, I build agents whose own thinking mechanisms evolve over time. The core idea is to treat long-horizon reasoning as an iterative control problem: agents should revise, critique, and improve their plans and judgments using explicit meta-level feedback or self-reflection on their thinking trajectories, rather than relying on static prompts or fixed reward models.



Our earlier work showed that explicitly modeling revision loops in writing reliably improves quality for both standalone agents and human–AI collaboration, highlighting structured test-time scaling as a principled mechanism for better outcomes [60, 40]. More recently, we focus on methods that update the agent’s internal policies and judges as it gains experience. Meta Policy Optimization (MPO) [60] reframes alignment as a two-level adaptive process by dynamically updating the reward model through meta-reward learning, reducing reward hacking by allowing evaluation criteria to change as training reveals new failure modes; *Metaⁿ* generalizes this recursion to multiple meta-levels of self-improvement [63]. Complementary inference-time approaches evolve meta-prompts during deployment, selectively updating them based on self-inconsistency signals to enable continual self-improvement without full retraining [40].

A key technical thread is converting test-time experience into reusable, evolvable knowledge. Instead of compressing long rollouts into a single prompt, we store per-instance reasoning experience as reusable thought-templates, then iteratively refine them using language feedback as text gradients, enabling long-context models to retrieve and apply these templates RAG-style across new tasks [36]. When tool-use experience introduces utility constraints such as budget, latency, and compute, we also model these signals explicitly and learn a meta-utility function at test time via bandit-style optimization, improving test-time decision making under constraints (ongoing work).

Our current emphasis is scalable synthetic collection and selection of complex workflow trajectories [74]. We developed an agent mid-training method, showing that agents fine-tuned with a small set of



high-quality trajectories substantially improve on frontier-discovery and agentic tasks. Across evolving agent frameworks (OpenEvolve, Terminus 2, and OpenHands), we find that trajectory quality and diversity dominate raw scale: selecting about 1K high-quality trajectories from roughly 200K—filtered for branching behavior, refinement strategies, and informative failures—can outperform conventional large-scale Best-of-N and generic test-time training approaches. Curated trajectory data functions as mid-training scaffolding that reduces train–test mismatch.

We explore meta-scaling frameworks (e.g., AsyncThink) that learn how to allocate branching and refinement budgets dynamically across stages of a workflow [24]. In domains with unverifiable judges and multi-step pipelines, such as data analysis and visual insight generation, fixed branching factors waste compute and miss opportunities. Our approach learns stage-specific scaling decisions, for example, increasing parallel exploration when outputs lack diversity or reallocating compute when judge scores are uninformative, yielding more strategic use of fixed compute budgets. Long term, I aim to unify human workflow modeling with self-evolving agent training into efficient agentic learning regimes: systems that learn from expert trajectories, accumulate test-time experience as structured memory, and continuously improve their reasoning policies and evaluation criteria.

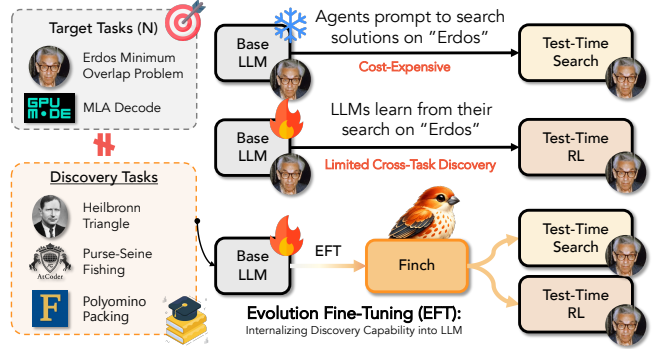


Figure 3: Evolution Fine-Tuning (EFT) [74]

2.6 Extension to Other Modalities

While my work primarily focuses on language, multidisciplinary collaborations have extended it to images and video, showing that multimodal models benefit from complementary cross-modal effects. Effective multimodal learning, however, requires careful design of fusion strategies and cognitive objectives such as reasoning and planning. Key results include the following: (i) vision–language models improve understanding of global connectivity and graph motif analysis in graph images [10], (ii) explicit temporal memory in video LLMs yields more coherent frames [4], (iii) aligning vision–language models with LLM feedback enhances reasoning in video and image tasks [53, 3, 2, 84], (iv) integrating multimodal representations into a unified latent space (a “Platonic” representation, where different modalities converge to a shared underlying structure) accelerates cross-modal understanding [35], and (v) LLMs can plan robotic manipulation tasks from language-based instructions [97]. Beyond multimodality, I also explore (i) autoregressive–diffusion hybrids for text generation [103], (ii) model uncertainty and calibrated abstention [38, 57, 66, 101] and generalization [1, 39], and (iii) theoretical analyses of learning frameworks [75].

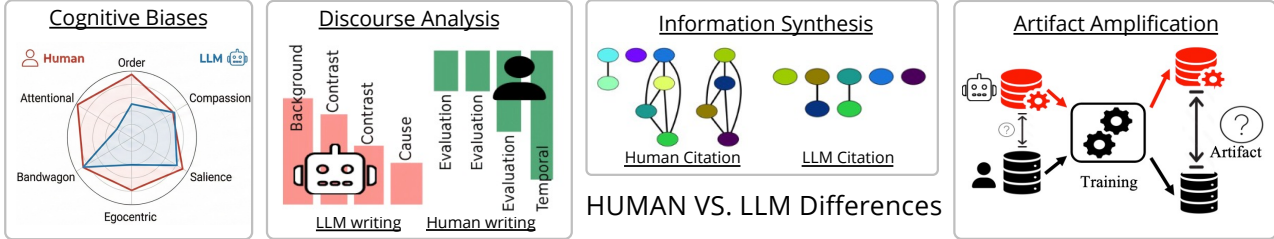
3 Cognitive Benchmarking and Evaluation

Evaluation is undergoing a structural shift as LLMs move from standalone models to persistent, tool-using collaborators embedded in real workflows. In this setting, leaderboard-style testing on fixed benchmarks is insufficient: evaluation becomes a measurement engineering problem, where the goal is to produce metrics that remain valid, interpretable, and safety-relevant as models, data, and deployment contexts drift; our recent benchmarks make evaluation principles and test difficulty explicit [33, 5]. My work develops evaluation methods and infrastructure across three directions: (i) quantifying the benefits and risks of AI-generated data and LLM-based evaluation, (ii) benchmarking human and machine cognition with theory-grounded assessments, and (iii) replacing static benchmarking

practices with long-horizon behavioral measurement for agentic systems [61, 82].

3.1 Potentials and Risks of AI-generated Data

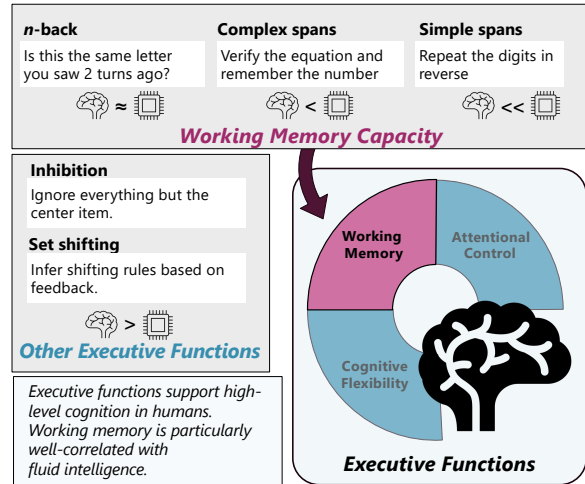
As LLMs began replacing or supplementing human annotation, I studied both the acceleration they enable and the new failure modes they introduce, including opaque error provenance and artifact reinforcement. On the opportunity side, I developed cost-reducing supervision and data-expansion techniques via symbolic rules [23, 52], annotation imputation [78], and information-theoretic measures [56, 57, 88, 66], as well as disagreement-aware sampling [95]. LLMs can generate annotations, prompts, simulated dialogues, and evaluation data, accelerating the research and development process. We have leveraged LLM-based evaluations [55, 3] and multi-agent simulations [43, 6, 98] to enrich training data.



In parallel, I characterized how overreliance can create an artificial data ecosystem (*Artificiality* [11]) that distorts measurement and behavior through cognitive biases (*CoBBLer* [67, 50]), discourse shifts [62], shallow synthesis [79], artifact amplification [11], and behavioral inconsistency across tasks and contexts [82, 94]. These results motivate evaluation protocols that explicitly model feedback loops between generators, judges, and downstream learners, rather than treating synthetic data and LLM judges as neutral measurement instruments.

3.2 Cognitive Assessment of LLMs

Traditional cognitive science experiments provide detailed insights into human cognition but are difficult to scale, creating a bottleneck for theory-driven data. To address this, we are developing *CogBench*¹, a benchmarking framework that connects cognitive psychology with artificial intelligence. *CogBench* integrates controlled human experiments with systematic evaluation of LLMs across executive control, working memory, temporal reasoning, and narrative comprehension. Although LLMs often perform well superficially, our findings reveal key cognitive limitations compared to humans. Models show a gap between capacity and control, with high working memory capacity but weak executive control and cognitive flexibility, as seen in tasks such as the Wisconsin Card Sorting Test [17]. They also rely heavily on prototypical patterns, producing inconsistent temporal and causal judgments in narratives [18, 91]. Moreover, while internal representations can detect incoherence, model outputs do not reliably distinguish coherent from incoherent narratives [16].



Limited executive functioning in LLMs [17] They also rely heavily on prototypical patterns, producing inconsistent temporal and causal judgments in narratives [18, 91]. Moreover, while internal representations can detect incoherence, model outputs do not reliably distinguish coherent from incoherent narratives [16].

By grounding AI evaluation in cognitive science and scaling it through benchmarking, *CogBench* characterizes the cognitive limits of current AI systems and provides the foundation for expert-facing systems where cognitive alignment can be assessed end to end.

Toward Long-horizon Agentic Evaluation. Finally, as LLMs become agents embedded in extended workflows, we move beyond static benchmarks toward behavioral measurement over long, multi-step episodes. Our recent studies show that current LLM agents are strong tool users but weak navigators

¹<https://minnesotanlp.github.io/CogBench-web/>

of long-horizon tasks [61], and that their behavior is often inconsistent across tasks and contexts [82], motivating persistent, workflow-embedded evaluation of agentic systems.

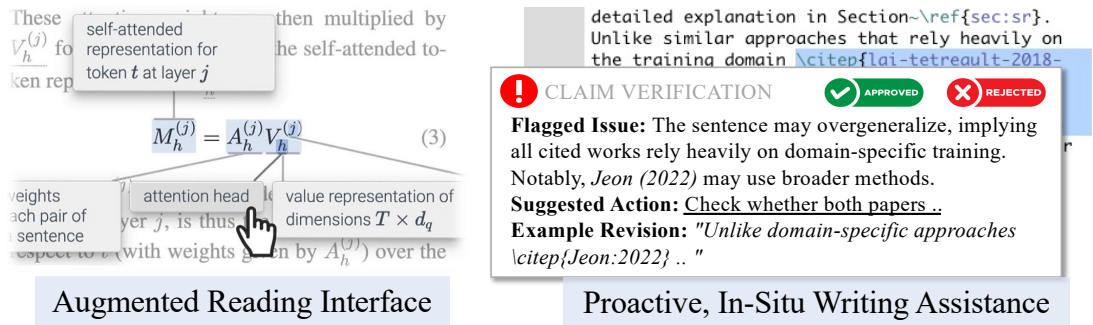
4 Thinking Assistants for Experts

In knowledge-intensive domains such as scientific research and legal practice, significant cognitive gaps remain between human experts and current LLM capabilities. These gaps arise from limitations in handling complex multi-step tasks, incorporating dynamic feedback, and adapting to evolving contexts. Full automation often oversimplifies the nuanced reasoning these tasks require, leaving AI systems misaligned with expert workflows.

My objectives are twofold: (i) benchmark AI cognition on end-to-end professional tasks to identify cognitive gaps, and (ii) create collaborative platforms where experts and AI agents interact, exchange feedback, and adapt over time. We focus first on building expert support systems in two flagship domains—*science* and *law*—and extend to media domains such as advertising and journalism (§4.3).

4.1 Assistants for Scientists

Reading Assistance. In collaboration with AI2 and UC Berkeley, I contributed to the development of *Semantic Reader* [34, 77], a tool publicly deployed on Semantic Scholar that provides in-situ definitions and explanations of terms as scientists read papers [47, 37], thereby reducing cognitive load. More recently, we launched *SciTalk*, which transforms research papers into short-form videos, making knowledge more accessible to wider audiences [87].



Writing Assistance. Scientific writing is a paradigmatic case of the iterative cognitive processes described in §1. I envision the future of scientific writing as a collaborative effort between scientists and AI agents, where both iteratively co-develop high-quality research outputs.

Writing assistance must go beyond supporting the immediate writing context. It should understand the full workflow, capturing the writer’s intent and providing in-situ support while minimizing distraction. I also focus on building interaction-centric AI systems that evolve through user feedback to better support scientific writing.

Our design principles include:

- *In-Situ Aids for Writing:* Document-level tools that detect overloaded symbols, redundant terminology, and logical inconsistencies using discourse-level modeling for real-time writing support.
- *Workflow-aware Interfaces:* Features that flag inconsistent scientific claims or overstated text and suggest revisions. For instance, a scientist drafting a related work section in Overleaf may trigger backend reasoning agents to automatically retrieve cited literature, detect errors or overstatements, and suggest improved citations with inferred rationales.
- *Human-Centric Design and Co-Learning:* Avoiding the “illusion of clarity” [86] by carefully designing interfaces that support scientists’ thinking processes while reinforcing human–AI collaboration.
- *Learning from Human Workflow:* Collecting human revision and writing workflow data to build cognitively-aligned systems for naturally integrated interactions [21, 71, 73, 98].
- *Learning from Interaction:* Adaptive systems that improve via user feedback, as demonstrated in collaborative editing [21] and iterative taxonomy building [72].

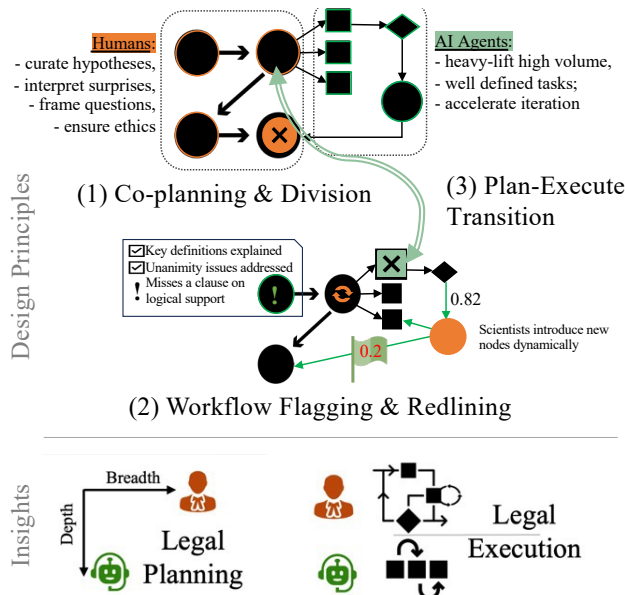
- *Trustworthy Models*: With NSF support, we develop foundation models integrating reliable, cross-modal knowledge for robust writing assistance.

Other Scientific Tasks. Beyond reading and writing, scientific workflows involve cognitively demanding stages such as peer review, collaboration, experimentation, and verification. Our lab is developing specialized agents and interfaces to support different stages of scientific tasks. For instance, we created *PeerRead* [42] (NAACL 2018), the first large-scale dataset of papers and reviews from computer science conferences, enabling research on acceptance prediction and aspect-specific review generation. In collaboration with visualization researchers, we are also building systems that automatically generate effective visualizations and insights (*A2P-Vis* [25]) and that scaffold the design of multi-agent analysis workflows (*FlowForge* [27]), adapting interactively to scientists’ needs.

4.2 Assistants for Lawyers

With Open Philanthropy and UMN Law School, we developed *LawFlow* [12], a benchmark for complex legal processing tasks requiring long-horizon reasoning and planning. The dataset captures detailed workflow data from both human lawyers and AI agents, including brainstorming, planning, research, and client communication for tasks such as advising a startup.

Our findings reveal that human legal reasoning is recursive and exploratory, whereas AI workflows are typically linear and exhaustive. From these human–AI comparisons, we derive design principles for legal interfaces, including collaborative planning, task division, and workflow flagging, and we are actively developing interfaces guided by these insights. Reflections from practicing lawyers [92] further highlight that generative AI already improves junior lawyer performance, yet adoption lags because evaluating AI outputs is cognitively demanding, especially for less experienced practitioners.



These challenges create several risks, including overtrust or overreliance on AI, shortcut-driven reasoning [70], and the potential intensification of expert polarization. Junior lawyers in particular risk being sidelined without proper scaffolding. To mitigate these risks, we emphasize the need for tools that upskill junior lawyers, train them to critically validate AI outputs, and provide workflow-aware assistance that supports lawyers as thinking and co-learning partners, ultimately lowering cognitive burden rather than replacing expertise.

4.3 Advertising and Journalism

Advertising and journalism are expert domains where nuanced affective responses, not just factual accuracy, drive outcomes. In collaboration

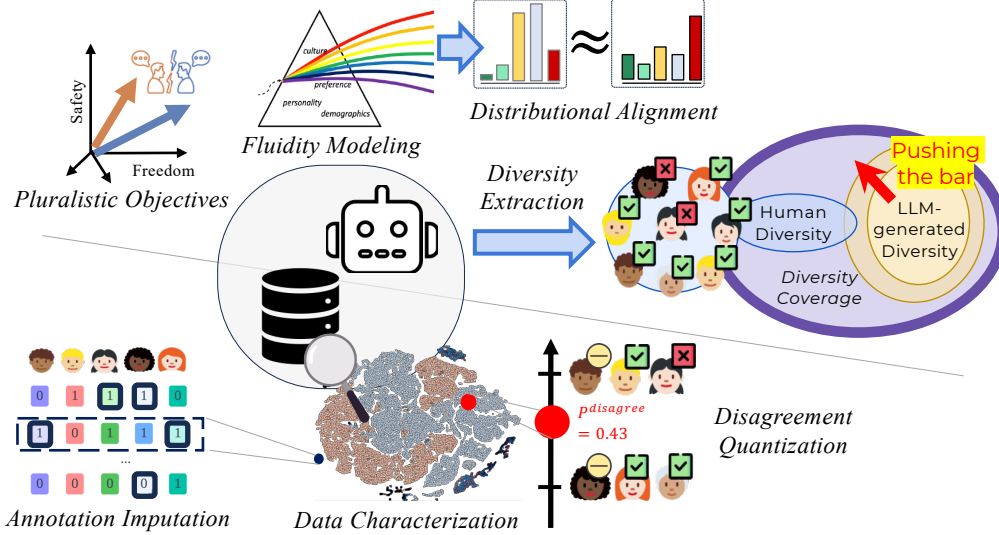
with Naver and MCAL, I contribute to developing computational metrics and algorithms for advertising in both academia and industry. For example, in the advertising market, we model nuanced affective responses such as skepticism [83] and nostalgia [58], moving beyond sentiment-based metrics. Recently, we have also studied how advertising is shifting from keyword-based search advertising toward generative ad platforms (e.g., Perplexity, Google AI Overviews), conducting comparative research on search patterns and consumer engagement to better understand the impact of AI on the advertising ecosystem [26]. Finally, I have worked on transforming the workflows of advertising services by integrating AI technologies. For instance, we are developing ad systems that support off-policy evaluation (OPE)—estimating the value of new ad policies from logged data without live deployment—in deterministic environments such as ad auctions [99].

By capturing detailed expert reasoning, we aim to build cognitively inspired models and expert-aware interfaces that adapt to the dynamic cognitive states of professionals. This approach will extend to other domains such as education [20], medicine, and programming, forming a broad *Expert Benchmarking* initiative to support specialized AI systems tailored to the cognitive demands of diverse expert roles.

5 Societal Alignment via Pluralism

Our society thrives on diverse perspectives shaped by different backgrounds, identities, and cultures, yet rapid AI advances often collapse them into a generalized “average.” Because cognition itself is pluralistic, alignment is not an add-on but a core requirement of cognitively-aligned AI: we must build people-centric technologies that capture and encourage diversity, ensuring inclusivity and societal alignment.

My goal is to build socially aware AI models that represent the full spectrum of human opinions through two complementary approaches: **data-centric** (§5.1) and **model-centric** (§5.2).



5.1 Data-centric Alignment

Diversity is difficult to define and capture: it manifests at individual, group, and community levels, and varies over time and context. Our “pluralistic representation” approach encodes disagreement and perspective variation directly into datasets. We develop methods to detect, characterize, and augment marginal viewpoints by (i) predicting missing annotations with collaborative filtering over annotations [78], (ii) analyzing annotation properties via learning dynamics [57], (iii) quantifying the disagreement level using demographic factors [95], and (iv) collecting perspectives from language learners to diversify model predictions [100, 54].

5.2 Model-centric Alignment

As NLP tasks grow more subjective, labels often shift from discrete to continuous values. Models that ignore this fluidity risk excluding certain groups. We design pluralistic alignment methods at multiple granularity levels: (i) *Individual preferences*: Model disagreement as a proxy for diversity using pairwise preference learning [55], (ii) *Group distributions*: Extract diverse opinions from LLMs [31] and align them with survey-based population distributions, and (iii) *Societal values*: Dynamically combine and resolve value conflicts using crowdsourced resolution scenarios [15].

As AI becomes embedded in education, industry, and politics, systems that ignore plurality risk amplifying monolithic perspectives and deepening polarization; pluralistic AI instead models coexistence among differing values, beliefs, and cultural backgrounds. To build a research community around these questions, I founded the Pluralistic Alignment workshop at NeurIPS 2024, bringing together ML, HCI, and social science researchers. My future work will operationalize pluralistic alignment through quantitative diversity metrics, disagreement-aware training objectives, and cross-cultural evaluation datasets.

References

- [1] Dyah Adila and Dongyeop Kang. Understanding out-of-distribution: A perspective of data dynamics. In *ICBINB@NeurIPS*, 2021.
- [2] Daechul Ahn, Yura Choi, San Kim, Youngjae Yu, Dongyeop Kang, and Jonghyun Choi. Isr-dpo: Aligning

large multimodal models for videos by iterative self-retrospective dpo. In *Thirty-Ninth AAAI Conference on Artificial Intelligence (AAAI)*, 2025.

- [3] Daechul Ahn, Yura Choi, Youngjae Yu, Dongyeop Kang, and Jonghyun Choi. Tuning large multimodal models for videos using reinforcement learning from ai feedback. In *Annual Meeting of the Association for Computational Linguistics (ACL)*, 2024.
- [4] Daechul Ahn, Daneul Kim, Gwangmo Song, Seung Hwan Kim, Honglak Lee, Dongyeop Kang, and Jonghyun Choi. Story visualization by online text augmentation with context memory. In *2023 IEEE/CVF International Conference on Computer Vision (ICCV)*. IEEE, October 2023.
- [5] Shady Ali and Dongyeop Kang. Who you test matters: Difficulty alignment in llm benchmark analysis. In *Workshop on Evaluating Evaluations (EvalEval) at ACL*, 2026.
- [6] Hyungjoo Chae, Yongho Song, Kai Tzu iunn Ong, Taeyoon Kwon, Minjin Kim, Youngjae Yu, Dongha Lee, Dongyeop Kang, and Jinyoung Yeo. Dialogue chain-of-thought distillation for commonsense-aware conversational agents. In *The 2023 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, 2023.
- [7] Seyyed Saeid Cheshmi, Hahnemann Ortiz, James Mooney, and Dongyeop Kang. Reasoning beyond literal: Cross-style multimodal reasoning for figurative language understanding. In *Findings of the Association for Computational Linguistics (EACL Findings)*, 2026.
- [8] Julia Christenson, Karin de Langis, Shirley Anugrah Hayati, and Dongyeop Kang. Effects of varying llm access on essay writing behavior. In *Workshop on Innovative Use of NLP for Building Educational Applications (BEA) at ACL*, 2026. Oral.
- [9] Debarati Das. *LMs as Sociotechnical Actors: A Scaffolding Framework for Recovering Context Lost in Text-Only Settings*. PhD dissertation, University of Minnesota, Minneapolis, MN, May 2026. Committee: Jaideep Srivastava, Dongyeop Kang, Stevie Chancellor, Jisu Huh.
- [10] Debarati Das, Ishaan Gupta, Jaideep Srivastava, and Dongyeop Kang. Which modality should I use - text, motif, or image? : Understanding graphs with large language models. In Kevin Duh, Helena Gomez, and Steven Bethard, editors, *Findings of the Association for Computational Linguistics: NAACL 2024*, pages 503–519, Mexico City, Mexico, June 2024. Association for Computational Linguistics.
- [11] Debarati Das, Karin De Langis, Anna Martin-Boyle, Jaehyung Kim, Minhwa Lee, Zae Myung Kim, Shirley Anugrah Hayati, Risako Owan, Bin Hu, Ritik Parkar, Ryan Koo, Jonginn Park, Aahan Tyagi, Libby Ferland, Sanjali Roy, Vincent Liu, and Dongyeop Kang. Under the surface: Tracking the artifactuality of llm-generated data, 2024.
- [12] Debarati Das, Khanh Chi Le, Ritik Sachin Parkar, Karin De Langis, Brendan Madson, Chad M. Berryman, Robin M. Willis, Daniel H. Moses, Brett McDonnell, Daniel Schwarcz, and Dongyeop Kang. LawFlow: Collecting and simulating lawyers’ thought processes. In *Conference on Language Modeling (COLM)*, 2025.
- [13] Debarati Das, David Ma, and Dongyeop Kang. Balancing effect of training dataset distribution of multiple styles for multi-style text transfer. In *Findings of the Annual Meeting of the Association for Computational Linguistics (ACL) Findings*, 2023.
- [14] Karin de Langis and Dongyeop Kang. A comparative study on textual saliency of styles from eye tracking, annotations, and language models. In Jing Jiang, David Reitter, and Shumin Deng, editors, *Proceedings of the 27th Conference on Computational Natural Language Learning (CoNLL)*, pages 108–121, Singapore, December 2023. Association for Computational Linguistics.
- [15] Karin de Langis, Ryan Koo, and Dongyeop Kang. Dynamic multi-reward weighting for multi-style controllable generation. In *Proceedings of the Conference on Empirical Methods in Natural Language Processing (EMNLP)*, 2024.
- [16] Karin de Langis, Püren Öncel, Ryan Peters, Andrew Elfenbein, Laura Kristen Allen, Andreas Schramm, and Dongyeop Kang. Mary, the cheeseburger-eating vegetarian: Do llms recognize incoherence in narratives? In *European Chapter of the Association for Computational Linguistics (EACL)*, 2026. Oral.
- [17] Karin de Langis, Jong Inn Park, Bin Hu, Khanh Chi Le, Andreas Schramm, Michael C. Mensink, Andrew Elfenbein, and Dongyeop Kang. Strong memory, weak control: An empirical study of executive functioning in llms. In *European Chapter of the Association for Computational Linguistics (EACL)*, 2026. Oral.

- [18] Karin de Langis, Jong Inn Park, Andreas Schramm, Bin Hu, Khanh Chi Le, and Dongyeop Kang. How llms comprehend temporal meaning in narratives: A case study in cognitive evaluation of llms. In *Proceedings of the Annual Meeting of the Association for Computational Linguistics (ACL)*, pages 29174–29191, 2025.
- [19] Karin de Langis, William Walker, Khanh Chi Le, and Dongyeop Kang. Tracing how annotators think: Augmenting preference judgments with reading processes. In *International Conference on Language Resources and Evaluation (LREC)*, 2026.
- [20] Nga Do, Jessica Wang, Zach Roper, Dongyeop Kang, and Richard Landers. Llm leadership coach for college students: Exploring the potential of llm-based chatbots for college leadership coaching, 2025. under review.
- [21] Wanyu Du, Zae Myung Kim, Vipul Raheja, Dhruv Kumar, and Dongyeop Kang. Read, revise, repeat: A system demonstration for human-in-the-loop iterative text revision. In *Proceedings of the First Workshop on Intelligent and Interactive Writing Assistants (In2Writing 2022)*, pages 96–108, Dublin, Ireland, May 2022. Association for Computational Linguistics.
- [22] Wanyu Du, Vipul Raheja, Dhruv Kumar, Zae Myung Kim, Melissa Lopez, and Dongyeop Kang. Understanding iterative revision from human-written text. In *Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 3573–3590, Dublin, Ireland, May 2022. Association for Computational Linguistics (ACL).
- [23] Steven Y. Feng, Varun Gangal, Dongyeop Kang, Teruko Mitamura, and Eduard Hovy. Genaug: Data augmentation for finetuning text generators. In *Deep Learning Inside Out (DeeLIO) Workshop at EMNLP 2020*, Online, November 2020.
- [24] Shuyu Gan, James Mooney, Pan Hao, Renxiang Wang, Mingyi Hong, Qianwen Wang, and Dongyeop Kang. Scaling unverifiable rewards: A case study on visual insights. In *Findings of the Association for Computational Linguistics (ACL Findings)*, 2026.
- [25] Shuyu Gan, Renxiang Wang, James Mooney, and Dongyeop Kang. A2p-vis: an analyzer-to-presenter agentic pipeline for visual insights generation and reporting. In *Workshop on Visualization for Generative AI (VISxGenAI) at IEEE VIS*, 2025. Honorable Mention.
- [26] Gabriel Garlough-Shah, Jong Inn Park, Shirley Anugrah Hayati, Dongyeop Kang, and Jisu Huh. Consumer engagement with ai-powered search engines: Implications for the future of search advertising research and practice. *Journal of Advertising Research*, 2026.
- [27] Pan Hao, Dongyeop Kang, Nicholas Hinds, and Qianwen Wang. Flowforge: Guiding the creation of multi-agent workflows with design space visualization as a thinking scaffold. In *IEEE Transactions on Visualization and Computer Graphics (IEEE VIS)*, 2025.
- [28] Shirley A. Hayati, Dongyeop Kang, Qingxiaoyang Zhu, Weiyan Shi, and Zhou Yu. Inspired: Toward sociable recommendation dialog systems. In *Conference on Empirical Methods in Natural Language Processing (EMNLP)*, Online, November 2020.
- [29] Shirley Anugrah Hayati, Taehee Jung, Tristan Bodding-Long, Sudipta Kar, Abhinav Sethy, Joo-Kyung Kim, and Dongyeop Kang. Chain-of-instructions: Compositional instruction tuning on large language models, 2024.
- [30] Shirley Anugrah Hayati, Dongyeop Kang, and Lyle Ungar. Does bert learn as humans perceive? understanding linguistic styles through lexica. *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, 2021.
- [31] Shirley Anugrah Hayati, Minhwa Lee, Dheeraj Rajagopal, and Dongyeop Kang. How far can we extract diverse perspectives from large language models? In *Proceedings of the Conference on Empirical Methods in Natural Language Processing (EMNLP)*, 2024.
- [32] Shirley Anugrah Hayati, Kyumin Park, Dheeraj Rajagopal, Lyle Ungar, and Dongyeop Kang. Stylex: Explaining styles with lexicon-based human perception. In *European Chapter of Association for Computational Linguistics (EACL)*, 2023.
- [33] Shirley Anugrah Hayati, Ruizi Wang, and Dongyeop Kang. From rubrics to recipe: Principle-centric benchmark for evaluating large language models. In *Workshop on Evaluating Evaluations (EvalEval) at ACL*, 2026.

- [34] Andrew Head, Kyle Lo, Dongyeop Kang, Raymond Fok, Sam Skjonsberg, Daniel S. Weld, and Marti A. Hearst. Augmenting scientific papers with just-in-time, position-sensitive definitions of terms and symbols. In *Proceedings of the SIGCHI Conference on Human Factors in Computing Systems (CHI)*, 2021.
- [35] Minoh Jeong, Min Namgung, Zae Myung Kim, Dongyeop Kang, Yao-Yi Chiang, and Alfred Hero. Anchors aweigh! sail for optimal unified multi-modal representations, 2024. under review.
- [36] Soyeong Jeong, Taehee Jung, Sung Ju Hwang, Joo-Kyung Kim, and Dongyeop Kang. When thoughts meet facts: Reusable reasoning for long-context lms. In *Findings of the Association for Computational Linguistics (ACL Findings)*, 2026.
- [37] Hwiyeol Jo, Dongyeop Kang, Andrew Head, and Marti A. Hearst. Modeling mathematical notation semantics in academic papers. In *Findings on Empirical Methods in Natural Language Processing (EMNLP Findings)*, 2021.
- [38] Taehee Jung, Dongyeop Kang, Hua Cheng, Lucas Mentch, and Thomas Schaaf. Posterior calibrated training on sentence classification tasks. In *2020 Annual Conference of the Association for Computational Linguistics (ACL)*, 2020.
- [39] Taehee Jung, Joo-Kyung Kim, Sungjin Lee, and Dongyeop Kang. Cluster-guided label generation in extreme multi-label classification. In *European Chapter of Association for Computational Linguistics (EACL)*, 2023.
- [40] Seungyeon Jwa, Daechul Ahn, Reokyoung Kim, Dongyeop Kang, and Jonghyun Choi. Becoming experienced judges: Selective test-time learning for evaluators. In *European Chapter of the Association for Computational Linguistics (EACL)*, 2026. Oral.
- [41] Dongyeop Kang. *Linguistically Informed Language Generation: A Multifaceted Approach*. PhD dissertation, Carnegie Mellon University, Pittsburgh, PA, May 2020. Available at https://kilthub.cmu.edu/articles/thesis/Linguistically_Informed_Language_Generation_A_Multi-faceted_Approach/24990648.
- [42] Dongyeop Kang, Waleed Ammar, Bhavana Dalvi, Madeleine van Zuylen, Sebastian Kohlmeier, Eduard Hovy, and Roy Schwartz. A dataset of peer reviews (PeerRead): Collection, insights and NLP applications. In *Meeting of the North American Chapter of the Association for Computational Linguistics (NAACL)*, New Orleans, USA, June 2018.
- [43] Dongyeop Kang, Anusha Balakrishnan, Pararth Shah, Paul Crook, Y-Lan Boureau, and Jason Weston. Recommendation as a communication game: Self-supervised bot-play for goal-oriented dialogue. In *Conference on Empirical Methods in Natural Language Processing (EMNLP)*, Hong Kong, November 2019.
- [44] Dongyeop Kang, Varun Gangal, and Eduard Hovy. (male, bachelor) and (female, ph.d) have different connotations: Parallely annotated stylistic language dataset with multiple personas. In *Conference on Empirical Methods in Natural Language Processing (EMNLP)*, Hong Kong, November 2019.
- [45] Dongyeop Kang, Varun Gangal, Ang Lu, Zheng Chen, and Eduard Hovy. Detecting and explaining causes from text for a time series event. In *Conference on Empirical Methods on Natural Language Processing (EMNLP)*, 2017.
- [46] Dongyeop Kang, Hiroaki Hayashi, Alan W Black, and Eduard Hovy. Linguistic versus latent relations for modeling coherent flow in paragraphs. In *Conference on Empirical Methods in Natural Language Processing (EMNLP)*, Hong Kong, November 2019.
- [47] Dongyeop Kang, Andrew Head, Risham Sidhu, Kyle Lo, Daniel Weld, and Marti A. Hearst. Document-level definition detection in scholarly documents: Existing models, error analyses, and future directions. In *Proceedings of the First Workshop on Scholarly Document Processing*, pages 196–206, Online, November 2020. Association for Computational Linguistics.
- [48] Dongyeop Kang and Eduard Hovy. Plan ahead: Self-supervised text planning for paragraph completion. In *Conference on Empirical Methods in Natural Language Processing (EMNLP)*, Online, November 2020.
- [49] Dongyeop Kang and Eduard Hovy. Style is NOT a single variable: Case studies for cross-stylistic language understanding. In *Proceedings of the 59th Annual Meeting of the Association for Computational Linguistics and the 11th International Joint Conference on Natural Language Processing (Volume 1: Long Papers)*, pages 2376–2387, Online, August 2021. Association for Computational Linguistics (ACL).

- [50] Dongyeop Kang*, Taehee Jung*, Lucas Mentch, and Eduard Hovy. Earlier isn't always better: Sub-aspect analysis on corpus and system biases in summarization. In *Conference on Empirical Methods in Natural Language Processing (EMNLP)*, Hong Kong, November 2019.
- [51] Dongyeop Kang, Tushar Khot, Ashish Sabharwal, and Peter Clark. Bridging knowledge gaps in neural entailment via symbolic models. In *Conference on Empirical Methods in Natural Language Processing (EMNLP)*, Brussels, Belgium, November 2018.
- [52] Dongyeop Kang, Tushar Khot, Ashish Sabharwal, and Eduard Hovy. Adventure: Adversarial training for textual entailment with knowledge-guided examples. In *The 56th Annual Meeting of the Association for Computational Linguistics (ACL)*, Melbourne, Australia, July 2018.
- [53] Jihyung Kil, Farideh Tavazoei, Dongyeop Kang, and Joo-Kyung Kim. Ii-mmr: Identifying and improving multi-modal multi-hop reasoning in visual question answering. In *Annual Meeting of the Association for Computational Linguistics (ACL) Findings*, 2024.
- [54] Haechan Kim, Junho Myung, Seoyoung Kim, Sungpah Lee, Dongyeop Kang, and Juho Kim. Learnervoice: A dataset of non-native english learners' spontaneous speech. In *Conference of the International Speech Communication Association (Interspeech)*, 2024.
- [55] Jaehyung Kim, Shin Jinwoo, and Dongyeop Kang. Prefer to classify: Improving text classifiers via auxiliary preference learning. In *International Conference on Machine Learning (ICML)*, 2023.
- [56] Jaehyung Kim, Dongyeop Kang, Sungsoo Ahn, and Jinwoo Shin. What makes better augmentation strategies? augment difficult but not too different. In *International Conference on Learning Representations (ICLR)*, 2022.
- [57] Jaehyung Kim, Yekyung Kim, Karin de Langis, and Dongyeop Kang. infoverse: A universal framework for dataset characterization with multidimensional meta-information. In *Proceedings of the Annual Meeting of the Association for Computational Linguistics (ACL)*, 2023.
- [58] Katie Kim, Bugil Chang, Debarati Das, Smitha Sudheendra, Jisu Huh, Dongyeop Kang, and Jaideep Srivastava. Rebuilding social connection and enhancing advertising effects through the nostalgic appeal during the pandemic. In *Association for Education in Journalism and Mass Communication (AEJMC)*, 2023.
- [59] Zae Myung Kim, Wanyu Du, Vipul Raheja, Dhruv Kumar, and Dongyeop Kang. Improving iterative text revision by learning where to edit from other revision tasks. In *Proceedings of the 2022 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, pages 9986–9999, Abu Dhabi, United Arab Emirates, December 2022. Association for Computational Linguistics.
- [60] Zae Myung Kim, Chanoo Kim, Vipul Raheja, and Dongyeop Kang. Toward evaluative thinking: Meta policy optimization with evolving reward models, 2025. under review.
- [61] Zae Myung Kim, Dongseok Lee, Jaehyung Kim, Vipul Raheja, and Dongyeop Kang. The amazing agent race: Strong tool users, weak navigators, 2026.
- [62] Zae Myung Kim, Kwang Hee Lee, Preston Zhu, Vipul Raheja, and Dongyeop Kang. Threads of subtlety: Detecting machine-generated texts through discourse motifs. In *Proceedings of the Annual Meeting of the Association for Computational Linguistics (ACL)*, 2024.
- [63] Zae Myung Kim, Young-Jun Lee, Seungyeon Jwa, and Dongyeop Kang. Metaⁿ: Recursive self-improvement through emergent depth, 2026. Under review for NeurIPS 2026.
- [64] Zae Myung Kim, Vassilina Nikoulina, Dongyeop Kang, Didier Schwab, and Laurent Besacier. Visualizing Cross-Lingual discourse relations in multilingual TED corpora. In *Proceedings of the 2nd Workshop on Computational Approaches to Discourse*, pages 165–170, Punta Cana, Dominican Republic and Online, November 2021. Association for Computational Linguistics.
- [65] Zae Myung Kim, Anand Ramachandran, Farideh Tavazoei, Joo-Kyung Kim, Oleg Rokhlenko, and Dongyeop Kang. Align to structure: Aligning large language models with structural information. In *Fortieth AAAI Conference on Artificial Intelligence (AAAI)*, 2026.
- [66] Ryan Koo, Yekyung Kim, Dongyeop Kang, and Jaehyung Kim. Meta-crafting: Improved detection of out-of-distributed texts via crafting metadata space (student abstract). In *Thirty-Eighth AAAI Conference on Artificial Intelligence (AAAI) Student Abstract*, 2024.

- [67] Ryan Koo, Minhwa Lee, Vipul Raheja, Jong Inn Park, Zae Myung Kim, and Dongyeop Kang. Benchmarking cognitive biases in large language models as evaluators. In *Annual Meeting of the Association for Computational Linguistics (ACL) Findings*, 2024.
- [68] Ryan Koo, Anna Martin, Linghe Wang, and Dongyeop Kang. Decoding the end-to-end writing trajectory in scholarly manuscripts. In *Proceedings of the Second Workshop on Intelligent and Interactive Writing Assistants (In2Writing) at CHI 2023*, 2023. https://minnesotanlp.github.io/REWARD_demo/.
- [69] Ryan Koo, Ian Yang, Vipul Raheja, Mingyi Hong, Kwang-Sung Jun, and Dongyeop Kang. Learning explainable dense reward shapes via bayesian optimization, 2025.
- [70] Nataliya Kosmyna, Eugene Hauptmann, Ye Tong Yuan, Jessica Situ, Xian-Hao Liao, Ashly Vivian Beresnitzky, Iris Braunstein, and Pattie Maes. Your brain on chatgpt: Accumulation of cognitive debt when using an ai assistant for essay writing task, 2025.
- [71] Khanh Chi Le, Linghe Wang, Minhwa Lee, Ross Volkov, Luan Tuyen Chau, and Dongyeop Kang. ScholaWrite: A dataset of end-to-end scholarly writing process. In *Annual Meeting of the Association for Computational Linguistics (ACL)*, 2026.
- [72] Minhwa Lee, Zae Myung Kim, Vivek Khetan, Alex Kass, and Dongyeop Kang. Cotaxon: Enhancing domain expertise through human-ai collaborative taxonomy construction, 2025. under review.
- [73] Minhwa Lee, Zae Myung Kim, Vivek A. Khetan, and Dongyeop Kang. Human-ai collaborative taxonomy construction: A case study in profession-specific writing assistants. In *CHI 2024 In2Writing Workshop*, 2024.
- [74] Young-Jun Lee, Seungone Kim, Minki Kang, Alistair Cheong Liang Chuen, Zerui Chen, Seung-ho Han, Taehee Jung, and Dongyeop Kang. Evolution fine-tuning: Learning to discover across 371 optimization tasks. In *Workshop on AI for Scientific Discovery in the Wild (AID-Wild) at CAIS*, 2026. Oral.
- [75] Chenliang Li, Siliang Zeng, Zeyi Liao, Jiayang Li, Dongyeop Kang, Alfredo Garcia, and Mingyi Hong. Joint reward and policy learning with demonstrations and human feedback improves alignment. In *International Conference on Learning Representations (ICLR)*, 2025. Spotlight.
- [76] Chu-Cheng Lin, Dongyeop Kang, Michael Gamon, Madian Khabsa, Ahmed Hassan Awadallah, and Patrick Pantel. Actionable email intent modeling with reparametrized rnn. In *Association for the Advancement of Artificial Intelligence (AAAI)*, 2018.
- [77] Kyle Lo, Joseph Chee Chang, Andrew Head, Jonathan Bragg, Amy X. Zhang, Cassidy Trier, Chloe Anastasiades, Tal August, Russell Authur, Danielle Bragg, Erin Bransom, Isabel Cachola, Stefan Candra, Yoganand Chandrasekhar, Yen-Sung Chen, Evie Yu-Yen Cheng, Yvonne Chou, Doug Downey, Rob Evans, Raymond Fok, Fangzhou Hu, Regan Huff, Dongyeop Kang, Tae Soo Kim, Rodney Kinney, Aniket Kittur, Hyeonsu Kang, Egor Klevak, Bailey Kuehl, Michael Langan, Matt Latzke, Jaron Lochner, Kelsey MacMillan, Eric Marsh, Tyler Murray, Aakanksha Naik, Ngoc-Uyen Nguyen, Srishti Palani, Soya Park, Caroline Paulic, Napol Rachatasumrit, Smita Rao, Paul Sayre, Zejiang Shen, Pao Siangliulue, Luca Soldaini, Huy Tran, Madeleine van Zuylen, Lucy Lu Wang, Christopher Wilhelm, Caroline Wu, Jiangjiang Yang, Angele Zamarron, Marti A. Hearst, and Daniel S. Weld. The semantic reader project: Augmenting scholarly documents through ai-powered interactive reading interfaces. In *Communications of ACM*, 2023.
- [78] London Lowmanstone, Ruyuan Wan, Risako Owan, Jaehyung Kim, and Dongyeop Kang. Annotation imputation to individualize predictions: Initial studies on distribution dynamics and model predictions. In *Workshop on Perspectivist Approaches to NLP (NLPerspectives) @ECAI 2023*, 2023.
- [79] Anna Martin-Boyle, Aahan Tyagi, Marti A. Hearst, and Dongyeop Kang. Shallow synthesis of knowledge in gpt-generated texts: A case study in automatic related work composition, 2024.
- [80] James Mooney, Zae Myung Kim, Young-Jun Lee, and Dongyeop Kang. Structure liberates: How constrained sensemaking produces more novel research output, 2026.
- [81] James Mooney, Vipul Raheja, Vivek Kulkarni, Ryan Koo, and Dongyeop Kang. Pairlens: Detecting correspondences in paired models, 2025. under review.
- [82] James Mooney, Josef Woldense, Zheng Robert Jia, Shirley Anugrah Hayati, My Ha Nguyen, Vipul Raheja, and Dongyeop Kang. Are llm agents behaviorally coherent? latent profiles for social simulation, 2025.

- [83] Smitha Muthya Sudheendra, Maral Abdollahi, Dongyeop Kang, Jisu Huh, and Jaideep Srivastava. SkOTaPA: A dataset for skepticism detection in online text after persuasion attempt. In Nicoletta Calzolari, Min-Yen Kan, Veronique Hoste, Alessandro Lenci, Sakriani Sakti, and Nianwen Xue, editors, *Proceedings of the 2024 Joint International Conference on Computational Linguistics, Language Resources and Evaluation (LREC-COLING 2024)*, pages 14871–14876, Torino, Italia, May 2024. ELRA and ICCL.
- [84] Jinwoo Nam, Daechul Ahn, Dongyeop Kang, Seong Jong Ha, and Jonghyun Choi. Zero-shot natural language video localization. *2021 IEEE/CVF International Conference on Computer Vision (ICCV)*, Oct 2021.
- [85] Rose Neis, Karin de Langis, Zae Myung Kim, and Dongyeop Kang. An analysis of reader engagement in literary fiction through eye tracking and linguistic features. In *Workshop on Narrative Understanding (WNU) @ACL 2023*, 2023.
- [86] C. Thi Nguyen. The seductions of clarity. *Royal Institute of Philosophy Supplements*, 89:227–255, 2021.
- [87] Jong Inn Park, Maanas Taneja, Qianwen Wang, and Dongyeop Kang. Stealing creator’s workflow: A creator-inspired agentic framework with iterative feedback loop for improved scientific short-form generation, 2025.
- [88] Ritik Sachin Parkar, Jaehyung Kim, Jong Inn Park, and Dongyeop Kang. Selectllm: Can llms select important instructions to annotate?, 2024.
- [89] Andy J. Phu, Karin de Langis, James Mooney, Khanh Chi Le, and Dongyeop Kang. SERUM: State extraction and refinement for user modeling. In *Conference on Language Modeling (COLM)*, 2026.
- [90] Vipul Raheja, Dhruv Kumar, Ryan Koo, and Dongyeop Kang. CoEdIT: Text editing by task-specific instruction tuning. In *The 2023 Conference on Empirical Methods in Natural Language Processing (EMNLP Findings)*, 2023.
- [91] Andreas Schramm, Karin de Langis, Jong Inn Park, Anh Thu Tong, Michael Mensink, Bin Hu, Khanh Chi Le, and Dongyeop Kang. How human is ai? a comparison of temporal meanings in 8 llms and 3 populations. In *American Association for Applied Linguistics (AAAL)*, 2025.
- [92] Daniel Schwarcz, Debarati Das, Dongyeop Kang, and Brett McDonnell. Thinking like a lawyer in the age of generative ai: Cognitive limits on ai adoption among lawyers. *Journal of Institutional and Theoretical Economics*, pages 6–26, 2026.
- [93] Zhecheng Sheng, Tianhao Zhang, Chen Jiang, and Dongyeop Kang. Bbscore: A brownian bridge based metric for assessing text coherence. In *Thirty-Eighth AAAI Conference on Artificial Intelligence (AAAI)*, 2024.
- [94] Ioannis Tsaknakis, Bingqing Song, Shuyu Gan, Jiangweizhi Peng, Dongyeop Kang, Alfredo Garcia, Gaowen Liu, Charles Fleming, and Mingyi Hong. Do llms recognize your latent preferences? a benchmark for latent information discovery in personalized interaction. In *Conference on Language Modeling (COLM)*, 2026.
- [95] Ruyuan Wan, Jaehyung Kim, and Dongyeop Kang. Everyone’s voice matters: Quantifying and distinguishing subjective annotation disagreement using demographic information. In *Thirty-Seventh AAAI Conference on Artificial Intelligence (AAAI)*, 2023.
- [96] Quan Wei, Chung-Yiu Yau, Hoi To Wai, Yang Zhao, Dongyeop Kang, Youngsuk Park, and Mingyi Hong. Roste: An efficient quantization-aware supervised fine-tuning approach for large language models. In *International Conference on Machine Learning (ICML)*, 2025.
- [97] Carl Winge, Adam Imdieke, Bahaa Aldeeb, Dongyeop Kang, and Karthik Desingh. Talk through it: End user directed manipulation learning. In *IEEE Robotics and Automation Letters (RA-L)*, 2024.
- [98] Ruixin Yang, Dheeraj Rajagopal, Shirley Anugrah Hayati, Bin Hu, and Dongyeop Kang. Confidence calibration and rationalization for LLMs via multi-agent deliberation. In *ICLR 2024 Workshop on Reliable and Responsible Foundation Models*, 2024.
- [99] Hongseon Yeom, Jaeyoul Shin, Soojin Min, Jeongmin Yoon, Seunghak Yu, and Dongyeop Kang. Breaking determinism: Stochastic modeling for reliable off-policy evaluation in ad auctions. In *ACM International Conference on Web Search and Data Mining (WSDM)*, 2026. Oral.

- [100] Haneul Yoo, Rifki Afina Putri, Changyoon Lee, Youngin Lee, So-Yeon Ahn, Dongyeop Kang, and Alice Oh. Rethinking annotation: Can language learners contribute? In *Annual Meeting of the Association for Computational Linguistics (ACL)*, 2023.
- [101] Skylar Zhai, Jingcheng Liang, and Dongyeop Kang. Abstain-r1: Calibrated abstention and post-refusal clarification via verifiable rl. In *Findings of the Association for Computational Linguistics (ACL Findings)*, 2026.
- [102] Tianhao Zhang, Zhecheng Sheng, Zhexiao Lin, Chen Jiang, and Dongyeop Kang. Bbscorev2: Learning time-evolution and latent alignment from stochastic representation. In *Conference on Empirical Methods in Natural Language Processing (EMNLP)*, 2025.
- [103] Hao Zou, Zae Myung Kim, and Dongyeop Kang. A survey of diffusion models in natural language processing, 2023.