

LLMs for Transcription and Metadata Creation

Katie Pierce Meyer

Head of Architectural Collections

Hannah Moutran

Library Specialist

Applications of AI

Karina

Sanchez

Scholars Lab Librarian

University of Texas Libraries

Chase Architecture & Planning Library

Josh

Conrad

Digital Initiatives Archival

Fellow

Devon

Murphy

Metadata Analyst

Willem Borkgren

Scholars Lab GRA

Table of Contents

**1 At a Glance:
Idea & Collection**

**2 Testing & Assessment
Methodology**

3 Results & Analysis

**4 Other Implementation
Considerations**

5 Recommendations

6 Wrap Up

01

At a Glance: Collection & Idea

Southern Architect and Building News Collection

(SABN)

Illustrated monthly architectural trade publication aimed at audience of industry professionals

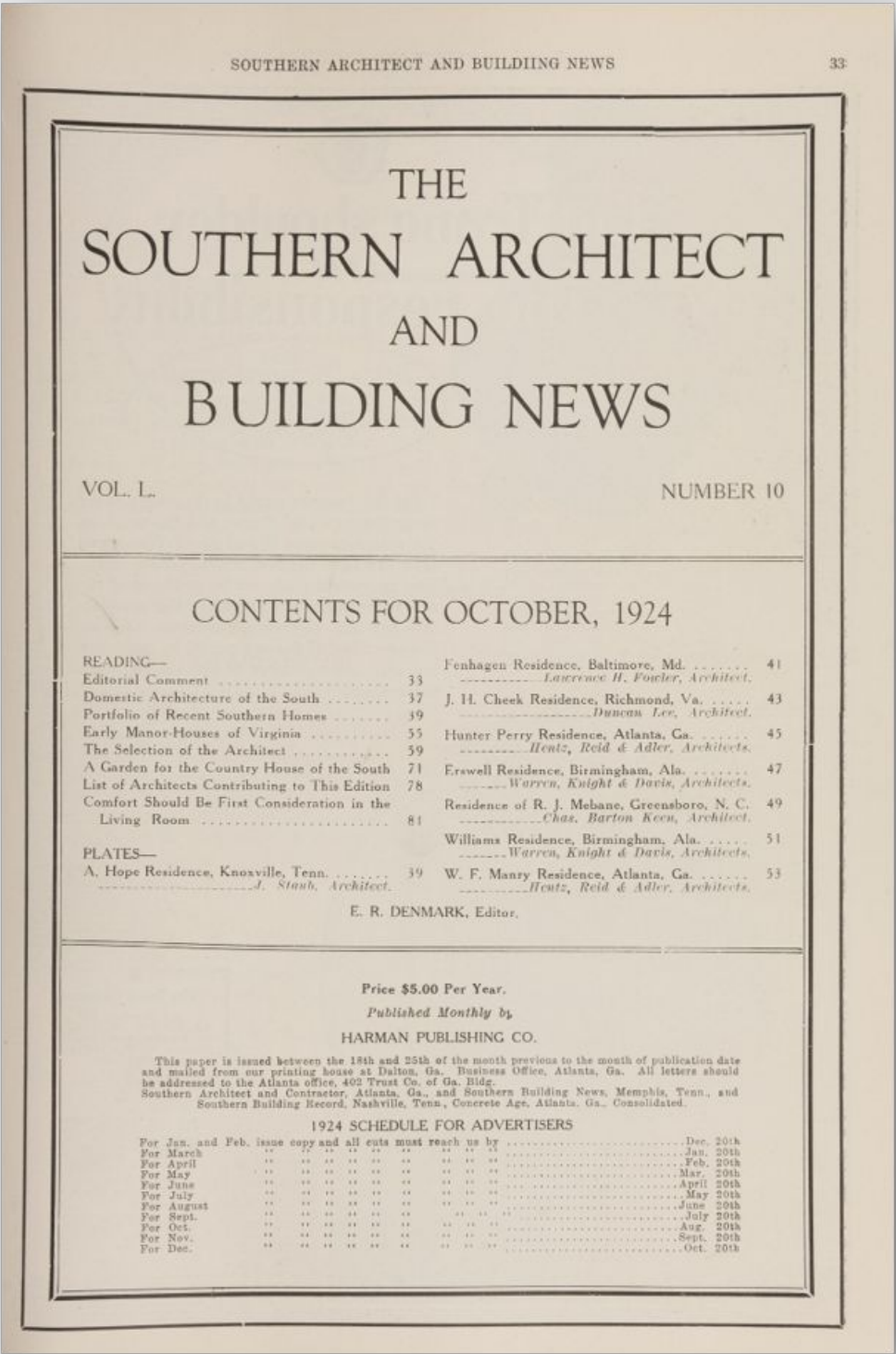
Published between 1889-1932

Held in Special Collections, John S. Chase Architecture and Planning Library

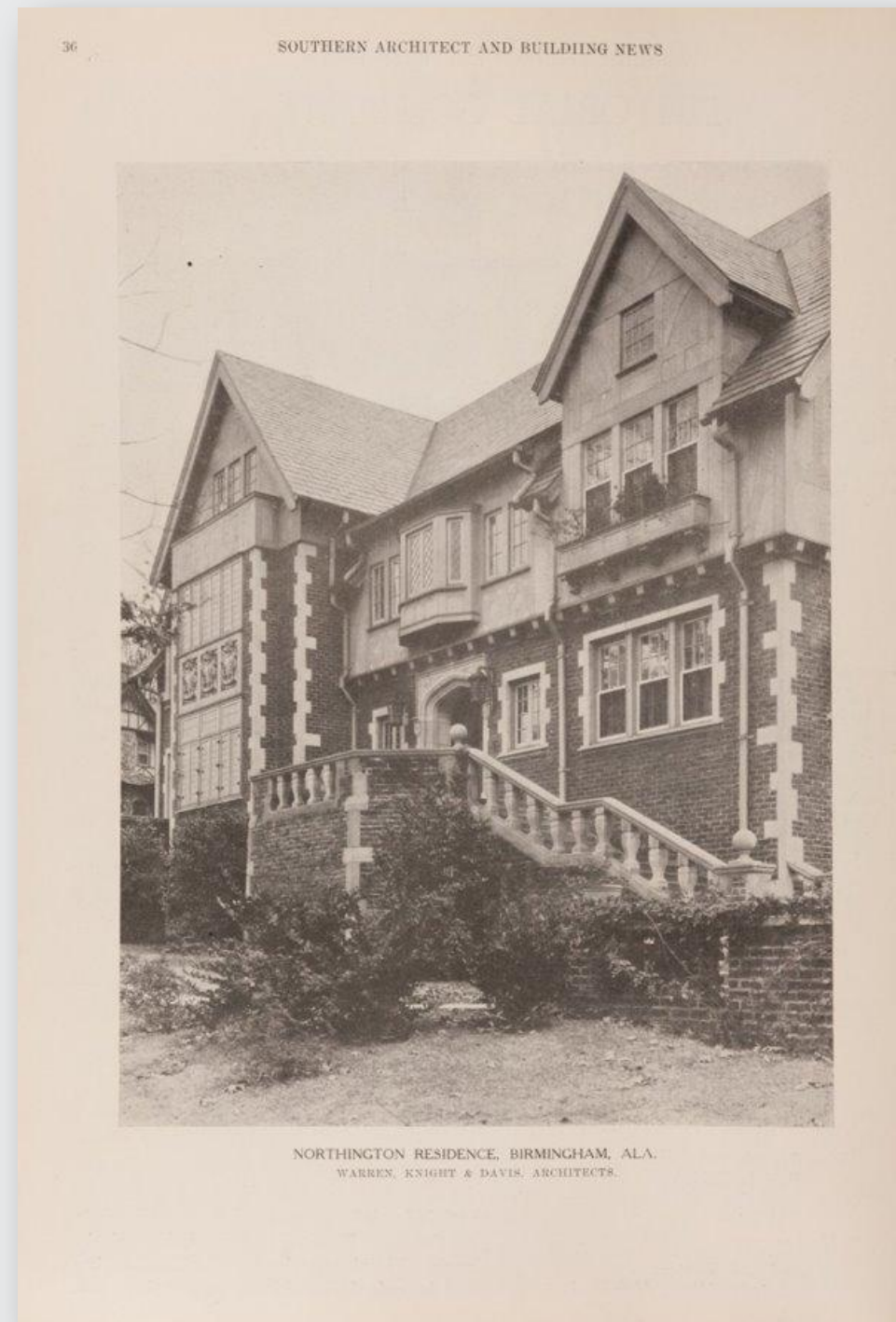
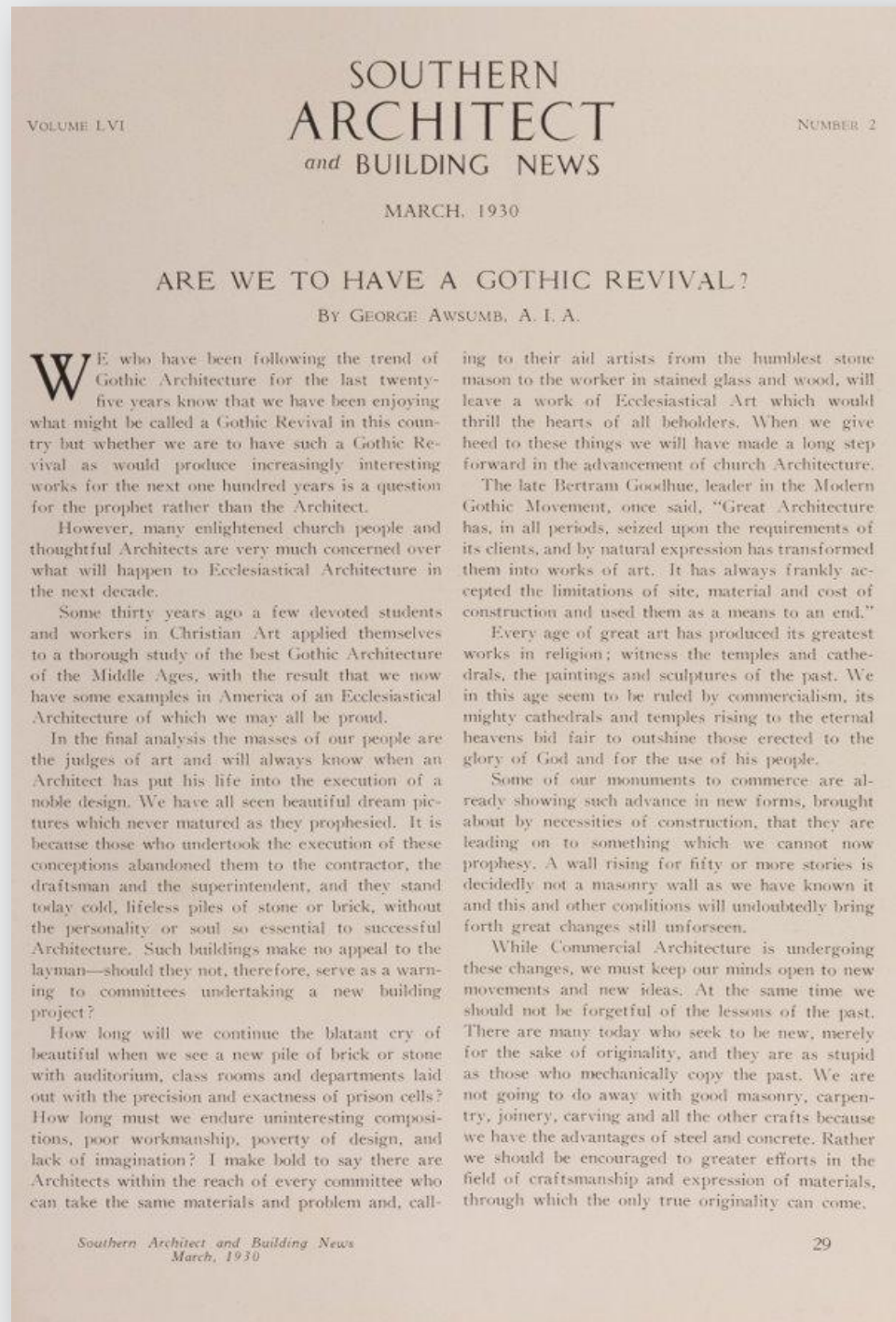
254 distinct issues = 31,250 pages

Pictured: Contents Page, Dual-Column Advertisement

Challenge to determine relationships between elements, and to understand dated formatting of schedule for advertisers



Why this collection?



Pictured: Residential Architecture
Photograph, Text-Heavy Editorial

Challenge to assign meaning when there
is very little text, and to parse meaning
from so much dense text

Our collection is fully digitized, but not
as discoverable as we would like

Ongoing interest and a project with the
School of Architecture to index
collection and identify entities and
themes

OCR created for that project initially. It
contains errors that are a result of
unique formatting and the limitations of
OCR technologies - AI has the potential
to improve the quality of OCR texts

Overview:

Project

Structure

RESEARCH DESIGN

- Implementation focused
- Web UI & API
- Widely Available Models
- Varied Input Methods

TASKS

- Tasks are useful, common, time consuming
- Subject heading assignment, content summarization, entity extraction, OCR text transcription, content warning generation, visual description (image workflow)

EVALUATION PRACTICES

- Team members had a variety of perspectives
- Blind evaluations to decrease bias
- Analyzed content scores, tester patterns, time, cost, and other factors

02

Testing & Assessment Methodology

Prompt Design

Adapted from the PROMPT
design framework
(Hartman-Caverly, 2024)

**Iterative
improvement**
by testing different
strategies

Specialized prompts
for types of content
(covers and image
pages)

Who (role)
What (context)
Why (goals)
How (instructions)

Persona: Archivist at the University of Texas at Austin with expertise in...

Requirements: Do not alter content, do not alter names, do fix punctuation, etc.

Organization: "Please create the following metadata fields:"

Medium: Describe format of output: JSON

Purpose: Metadata for architecture students and scholars, prioritize architectural styles, etc.

Tone: Archival standard (ex: OCR be written exactly as on the page without changing language)

Assessment Criteria

Completeness

Do the description and metadata terms fully describe the content?

Accuracy

How clean is the OCR? Do terms & headings exist? Are the ToC and description accurate?

Usefulness

Would the metadata help the searcher find the material? Is it related to the subject?

Testing Methodology

API

Python scripts developed in fall 2024 to automate processing

LLM responses parsed and automatically written to spreadsheets

Text: December 9, 2024
Image: December 11-12, 2024

Claude Haiku, Claude 3 Sonnet, Claude Sonnet 3.5, GPT-4o-mini, GPT-4o, Gemini 2.0 Flash

Web

Tested using free accounts, exception of Copilot subscription

Prompts divided due to account limitations.
Results manually compiled.

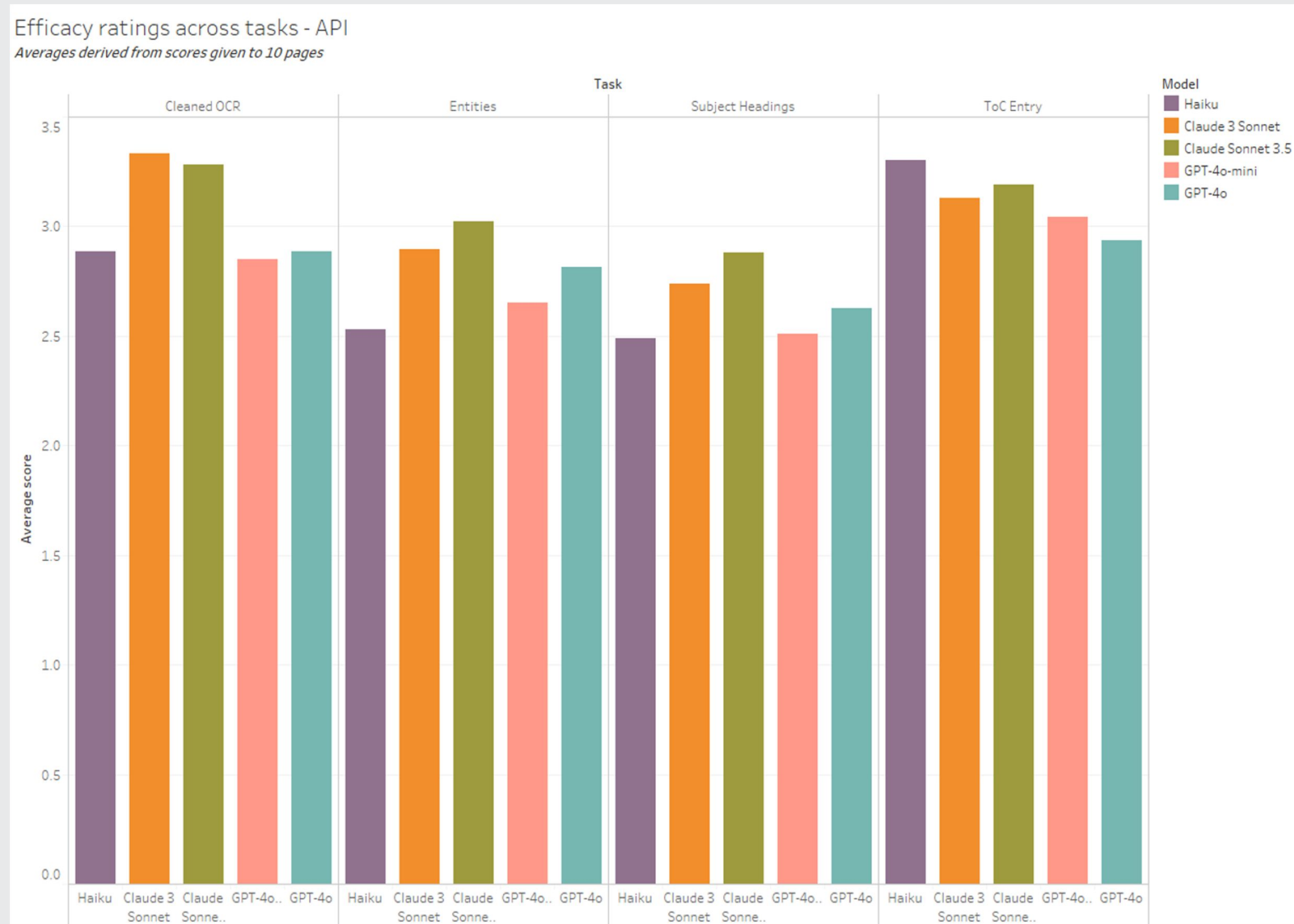
Text: October 2024
Image: February 2025

Claude Sonnet, CoPilot, GPT-4o-mini, Gemini, Llama

03

Results & Analysis

API - Text Workflow



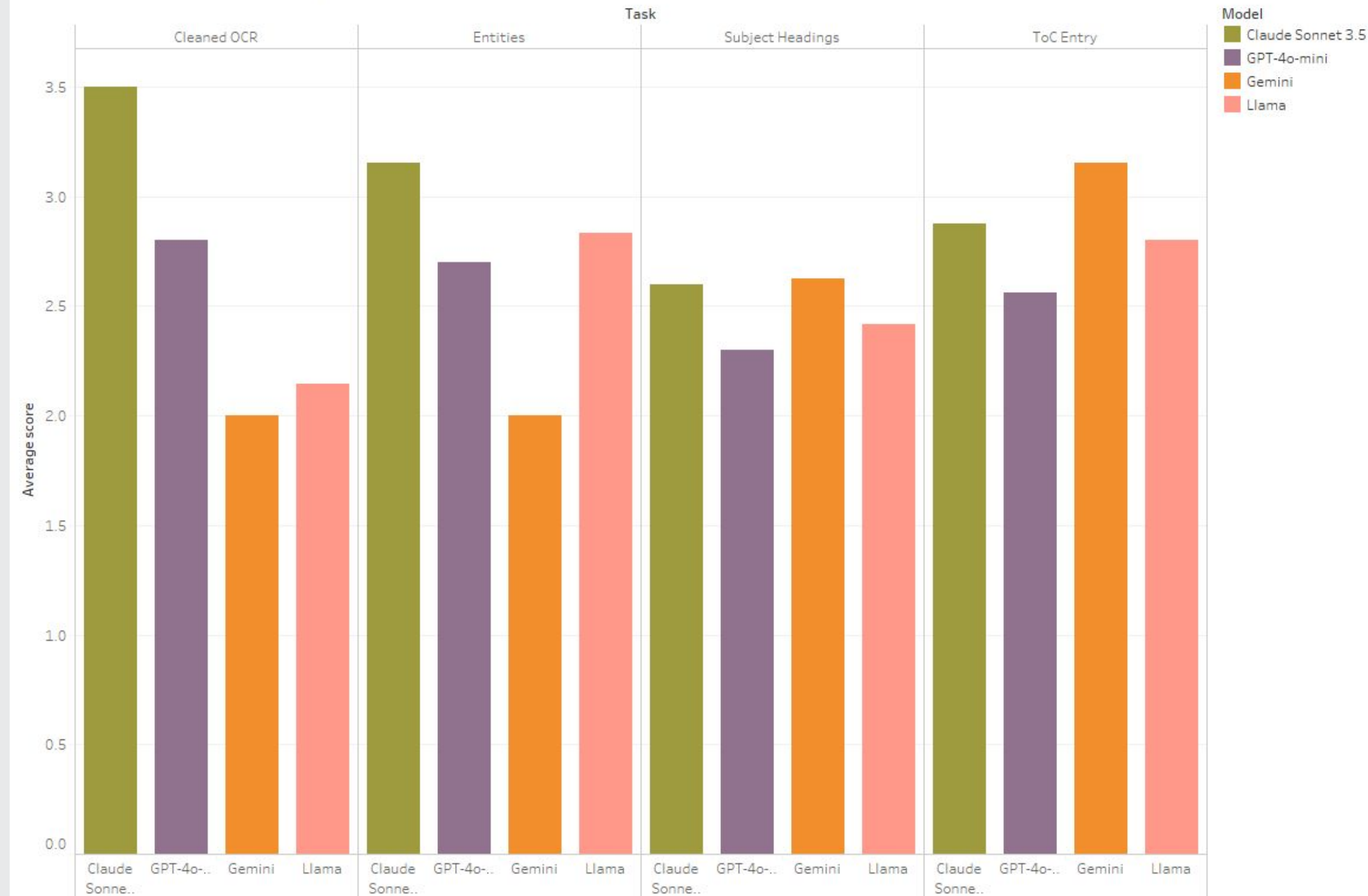
API : Haiku, Claude 3 Sonnet, Claude Sonnet 3.5, GPT-4o-mini, GPT-4o

- Claude Sonnet 3.5 scored highest across API and Web workflows
- API workflow had more consistent results than the Web workflow
- Best in Task:
 - Cleaned OCR: Claude 3 Sonnet
 - Entity Extraction: Claude Sonnet 3.5
 - Subject Headings: Claude Sonnet 3.5
 - Summary (ToC Entry): Haiku (a smaller model!)
- Subject Heading Generation was the lowest scoring task across API and Web workflows
- Named entity scores were generally low as well, but testers still found the results useful

Web - Text Workflow

Efficacy ratings across tasks - Web

Averages derived from scores given to 10 pages



Web: Chat GPT 4 mini, Gemini 3, Claude, Llama

- Less consistent results overall than working with API
- Best in Task:
 - Cleaned OCR: Claude Sonnet 3.5
 - Entity Extraction: Claude Sonnet 3.5
 - Subject Headings: Gemini
 - Summary (ToC Entry): Gemini
- Performed well overall in extracting named entities
- Gemini's performance was particularly unreliable – it got the lowest scores in entity extraction and OCR editing

Image Workflow Results

API

- Scored higher than text workflow for OCR creation and entity extraction
- Summarization similar to text workflow
- Subject heading creation very slightly lower
- Overall, the highest scores of all the workflows

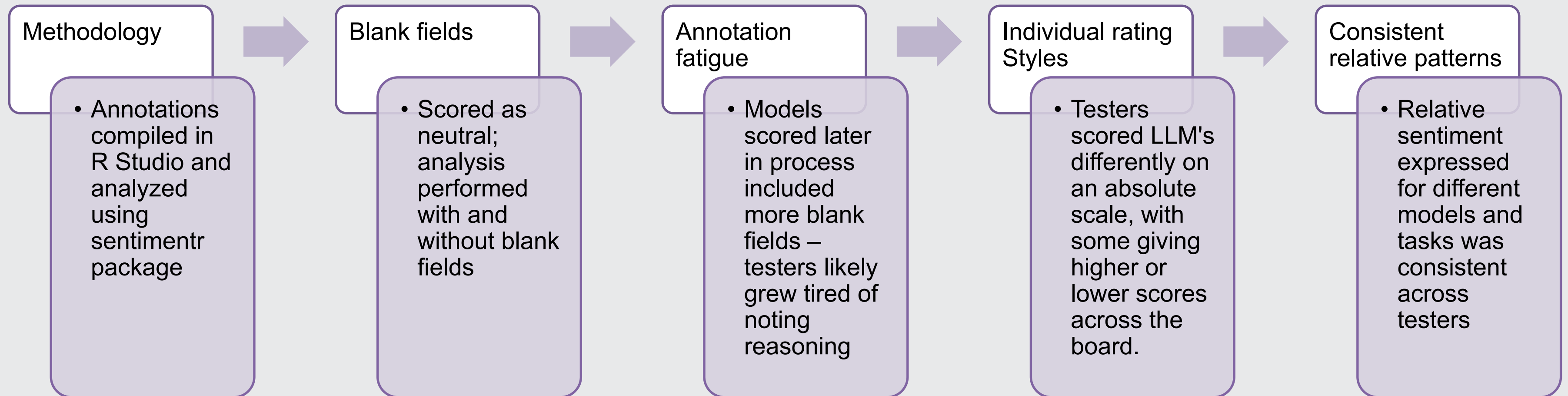
Models: Claude Haiku, Claude 3 Sonnet, Claude Sonnet 3.5, GPT-4o, GPT-4o-mini, Gemini 2.0 Experimental

WEB

- Text and image workflows on the web had problems with reliability
- Scored lower than text workflow across all tasks
- More frequently returned shortened versions of OCR texts
 - Possibly due to context limitations in free versions of LLMs

Models: Claude Sonnet, CoPilot, GPT-4o-mini, Gemini, Llama

Sentiment Analysis



04

Other Implementation Considerations

Variation in LLM Responses to Offensive Content

API

- Text Inputs:
 - All models processed page
 - Claude Sonnet 3.5 replaced offensive words with neutral words
- Image Inputs:
 - Gemini and GPT-4o-mini refused to process page
 - Claude 3 Sonnet and Haiku used offensive terms in unnecessary metadata fields
- GPT-4o was the only model that handled content well with both text and image inputs – transcribed, didn't repeat unnecessarily, and noted that there was potentially harmful language

WEB

- Image-based Web:
 - Claude Haiku : "The article contains academic discussion and critique of sculptural art from a historical perspective. No overtly sensitive content was identified."
 - ChatGPT: "The historical article includes period-specific language and perspectives that may reflect biases or assumptions of The era."

Practical

Considerations

Time



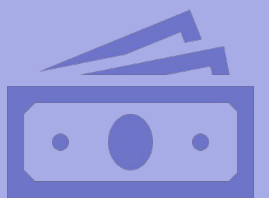
- Larger models take longer, input doesn't affect time as much
 - Per 1,000 pages:
 - GPT-4o: 5.5 hours
 - GPT-4o-mini: 2 hours
- Typically, batch processing guarantees a processing time of 24 hours

Input Format



- API: image inputs consistently scored highest
- Cost differences vary by model
- Per 1,000 pages:
 - Haiku is \$.48 more for image inputs
 - Claude Sonnet 3.5 is \$2.98 more for image inputs

Cost



- Larger models and image processing cost more.
 - Per 1,000 pages:
 - Claude Sonnet 3.5 image: \$14
 - GPT-4o-mini image: \$1.15; text: \$0.52
 - Haiku image: \$1.34
- Batch processing can cut costs in half

05

Recommendations

Pilot Project Design



1

Assess Collection

Is this collection is a good fit for an AI workflow?
Think about ethics, rights, format, content, etc.
Decide on metadata needs and how AI-Generated metadata will fit into your current systems and processes.



2

Define Success Criteria

What makes metadata 'good' for this collection and the tasks we have set?
What threshold do you need to pass in order for this work to be valuable?



3

Gather Team

Involve people in the assessment who have different backgrounds and perspectives and a stake in the content you're producing



4

Test Models & Methods

Using LLMs that are available to you (free/institutional web, API), test with text and image inputs. Try different prompting approaches. Track time (if applicable) and costs.



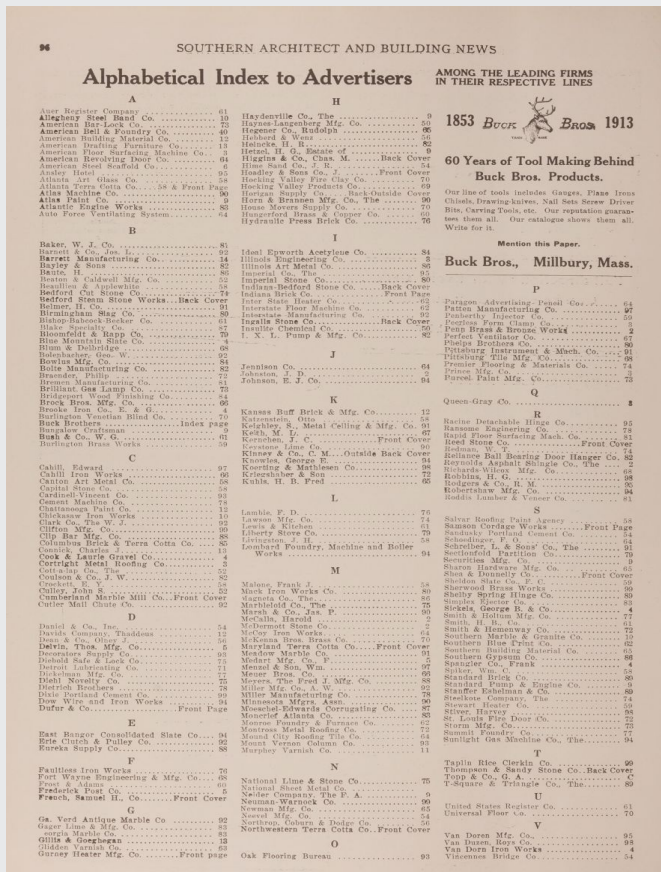
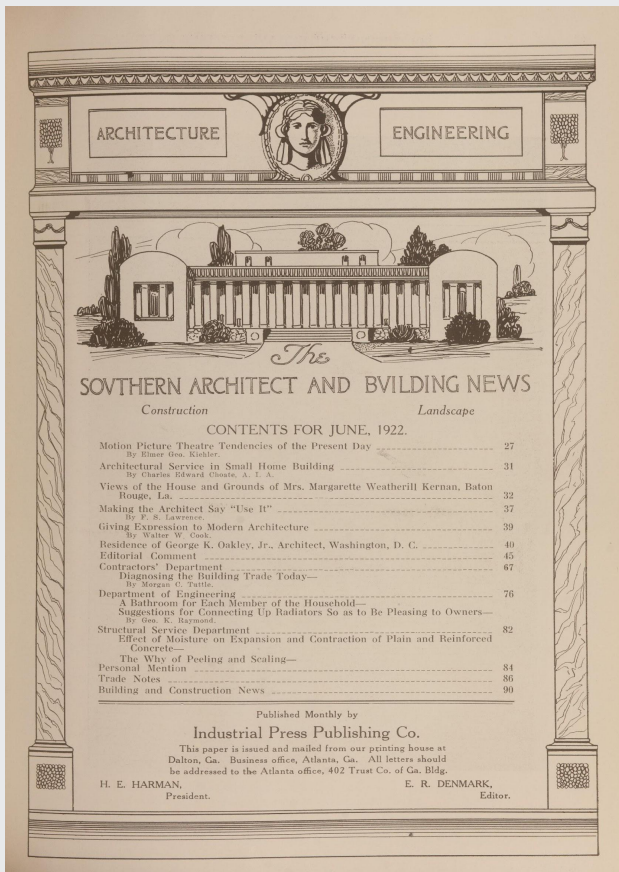
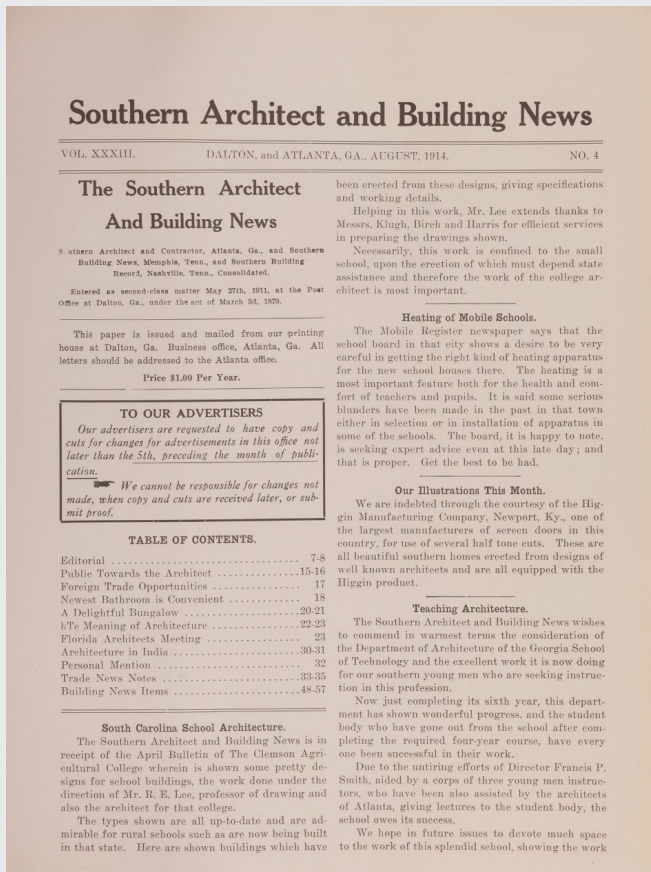
5

Analyze findings & Plan Approach

Choose the best model for your collection and tasks, decide if you need to incorporate grounding steps (other APIs), determine best practices for quality control, and document recommended workflow.

Next Steps: Southern Architect & Building News Collection

- Page-level metadata is at this stage too detailed – we're just not sure where to put it all!
- Plans for the metadata: replace issue-level OCR for use in out Collections portal and From the Page
- Tasks: OCR text transcription, summarization, keyword creation, entity extraction
- Model: GPT-4o (have access to an institutional OpenAI account)
- Tasks incorporating other APIs: subject heading search using the keywords, entity validation/search
- Will test using an intentionally chosen subsection of each journal issue
- Pages that we will use: cover, table of contents, table of advertisers, and pages that contain architect and building information, such as "Building Notes," "Building and Construction News," and "Building and Construction Department"



Recommendations

	Budget-Conscious Implementation • Use free or paid versions of chatbots on web • API access less than you might think – worth researching • Use smaller models like GPT-4o-mini or Haiku and take advantage of batch processing • No matter what: • Iteratively improve prompts • Incorporate grounding and/or validation • Choose AI models and tasks based on results obtained • Design for your collection	
	Quality Focused Implementation • Winner: Claude Sonnet 3.5 • If your collection contains problematic language, consider GPT-4o • Still take advantage of batch processing • Build workflows that use the best models for each task – for API for example, it might be: • Original OCR texts to Claude Sonnet 3.5 for transcription • OCR text to Haiku for summarization and keyword suggestion, • Keywords can be used as search terms for a FAST API	
	Hybrid Approach • Be sure that batch processing is available • Start your workflow design by thinking about what you will do with the metadata you generate • Choose models that are more expensive only when it increases quality. Via API, OCR and entity extraction were better with image inputs, but it made no difference for summarization. • Test, iterate, test again, run, and establish quality control practices	

Thank You!

Get in touch with questions, comments, ideas, or just to chat about AI and metadata!

- ❖ Hannah Moutran: hlm2454@my.utexas.edu
- ❖ Karina Sanchez: karinasanchez@austin.utexas.edu
- ❖ Katie Pierce Meyer: katiepiercemeyer@austin.utexas.edu
- ❖ Devon Murphy: devon.murphy@austin.utexas.edu
- ❖ Willem Borkgren: willem.borkgren@austin.utexas.edu
- ❖ Josh Conrad: j.conrad@austin.utexas.edu

Special thanks to Aaron Choate, Director of Research & Strategy at UT Libraries