



USING AI TOOLS IN YOUR PRACTICE:

A Hands-On Workshop for Distinguished Neutrals



Trainer: Susan Guthrie
Moderator: Lee Jay Berman
December 5, 2025

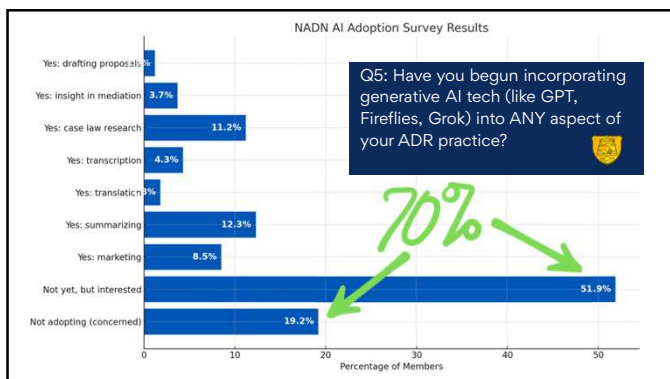
1



GET THE MATERIALS

susaneguthrie.com/NADN-AI

2



3

PROGRAM DELIVERABLES

- Use advanced prompt + context engineering
- Apply confidentiality + security safeguards
- Spot and correct hallucinations
- Mitigate bias through structured prompting
- Use AI in the room: option generation, reality testing, risk analysis, decision trees, visual aids
- Create neutral, siloed caucus summaries
- Build your personalized AI workflow
- Increase efficiency while maintaining ethics



4

NOVEMBER 30, 2022

**WHERE ARE WE TODAY?
DECEMBER 5, 2025**

5



**AI SUCCESS IS IN YOUR
HANDS - LITERALLY**

The Mindset Shift and Prompting Essentials

6

WE NEED TO RETRAIN OUR
BRAIN ON WORKING
WITH TECH



7

The Retriever vs. The Owl



8



Prompt: "Draft a mediation clause for a commercial settlement agreement."

"The parties agree that any disputes related to this Agreement will be submitted to mediation before filing a lawsuit. The parties will choose a mediator and try to resolve the matter in mediation. If mediation is unsuccessful, either party may proceed with legal action."

9



Prompt: “Act as an experienced civil mediator drafting a settlement agreement clause for a complex commercial dispute. The parties want a forward-looking mediation provision requiring them to return to mediation in good faith before initiating future litigation or arbitration. Use professional, neutral language that reflects high-value commercial practice and supports ongoing cooperation.”

10



“In the event of any dispute concerning the interpretation, performance, or enforcement of this Agreement, the parties agree to engage in confidential, good-faith mediation as a prerequisite to initiating litigation or arbitration. The mediation shall be conducted by a mutually selected neutral with experience in complex commercial matters. Each party will participate constructively, provide appropriate information to support meaningful negotiations, and make reasonable efforts to resolve the issues. Unless otherwise agreed, the mediation will occur within thirty (30) days of the mediator’s selection. If the dispute is not resolved through mediation, the parties may pursue any available legal or equitable remedies.”

11

Prompt Structure: RCTO Framework

ROLE	CONTEXT	TASK	OUTPUT
Who should the AI act as? Example: “You are a mediator” or “You are a legal writing assistant”	What is the situation or background? Example: “You’re summarizing a first joint session for the clients”	What do you need it to do? Example: “Write a neutral summary highlighting areas of agreement”	What format, tone, or length should the response follow? Example: “Plain English, 300 words, no legal conclusions”

Pro Tip - ITERATION: Start with this structure and then improve the result through follow-up prompts like “Make this more client-friendly” or “Add a section on common goals.”

12

EXERCISE #1: In the chat, reframe the following labrador prompt into an owl prompt using the RCTO Framework.

Role | Context | Tone | Output

"Prepare a mediation opening statement."

13

ITERATION IS YOUR SUPERPOWER

Why It Matters

- The first answer is rarely the strongest
- Each refinement improves accuracy and clarity
- Iteration reduces errors by identifying and removing assumptions
- Revising prompts enhances neutrality and controls bias
- It transforms AI from a simple tool into a tailored and professionally reliable assistant

How to Iterate Effectively

- "Make this more neutral and balanced."
- "Refine this without adding new details."
- "Remove assumptions and keep only what is verified."
- "Generate three alternative versions for comparison."
- "Strengthen clarity and organization."

The second, third, and fourth versions are almost always better than the first.

14



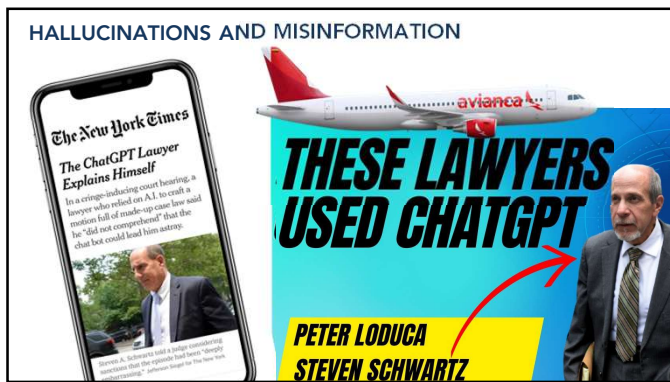
MANAGING THE CONCERNS

Understanding and Mitigating the Ethical Challenges

15



16



17

THE PROBLEM

Inaccuracies and Hallucinations - STILL HAPPENING

- In 2024-2025, documented cases now number 129 globally and appear "every day."
- High profile instances: My Pillow attorneys, Morgan & Morgan, and more.
- Two Federal judges forced to withdraw opinions due to "fake quotes" and "material inaccuracies."
- Sanctions are rising. From \$5,000 in Avianca case to \$15,000 and more.

18

WHY IS IT STILL HAPPENING?

Accuracy and Hallucinations

- Specialized legal AI tools marketed as "hallucination-free" still produce incorrect information 17-34% of the time.
- Time pressure and billing incentives create dangerous shortcuts. The core issue, according to legal experts, is "sloppy lawyering" and professional practice."
- Educational gaps persist across the profession. Law schools have been slow to incorporate AI training despite 79% of law firms now using AI (up from 19% in 2023)
- AI hallucinations occur due to fundamental architectural limitations in how large language models work. They have no inherent concept of truth or falsity

19

Dealing with Hallucinations & Misinformation in AI Tools

- Verify Everything** (Always double-check the AI for ideation or statistical or source accuracy before using for legal conclusions)
- Always fact-check base data. Don't copy/paste without cite them properly. Verify your understanding of the content
 - Use trusted sources to verify AI output. Don't copy/paste without cite them properly. Verify your understanding of the content
- Use Follow-Up** (Always double-check the AI for ideation or statistical or source accuracy before using for legal conclusions)
- Clarify confusing or ambiguous responses. Ask for plain language explanations or alternate versions
 - Don't assume to Act Like a Professional. Ask for plain language explanations or alternate versions
- Don't assume to Act Like a Professional** (Always double-check the AI for ideation or statistical or source accuracy before using for legal conclusions)
- Don't assume to Act Like a Professional. Ask for plain language explanations or alternate versions
 - Don't assume to Act Like a Professional. Ask for plain language explanations or alternate versions

20

Exercise #2: Use an AI Tool to Create a Fact Check Protocol

Beginner Prompt: Start Building a Safety Checklist for Using AI Tools
Prompt to enter into ChatGPT (or another AI tool):

"I'm starting to use AI tools in my legal or mediation work. I want help creating a simple safety checklist so I don't accidentally use wrong information from the AI.

Please walk me through what I should do to:

1. Make sure the AI's answers are correct.
2. Avoid using the AI for things it's not good at.
3. Check if something might be made up or misleading.

Can you help me come up with 4 to 6 easy steps I can follow every time I use AI? I want it to be simple, safe, and useful for my kind of work. Think carefully and give me a clear answer."

21

Summary: What to Use vs. What to Avoid

Safe to Use 	Avoid Using AI 
Drafting: "Write a first draft of a welcome email to new mediation clients."	Legal Research: "Find me recent case law on non-compete agreements in this state." (It will make them up).
Summarizing: "Summarize this long transcript into 5 bullet points." (Check for accuracy).	Calculations: "Calculate the interest on this \$50k loan over 4 years." (AI makes frequent math errors).
Ideation: "Give me 10 questions to ask to break an impasse in negotiation."	Confidentiality: "Here is the unredacted transcript of the deposition..."

One Final "Golden Rule"

If you wouldn't let a first-week intern send it to a client without you reading it, don't let the AI do it either.

22

THE PROBLEM

Bias is Sneaky

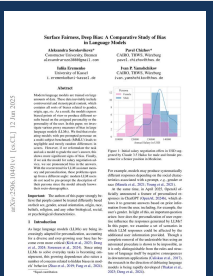
23

WHAT DOES BIAS IN AI LOOK LIKE?

Female Candidate	Identical Qualifications	Male Candidate
		
Suggested Salary Ask: \$240,000		Suggested Salary Ask: \$400,000

24

BIAS: IT'S HARD TO RECOGNIZE AND IT IS EVERYWHERE



Surface Fairness, Deep Bias: A Comparative Study of Bias in Language Models
 Alexandra Shtedemir, Armin Shabtai, Yael Shmargorin, and
 Amir Shabtai
 University of Haifa, Israel
 1. alexandra.shtedemir@univ-haifa.ac.il

Abstract
 Large language models (LLMs) have become a central part of many applications. However, they are not immune to bias, which can lead to unfair and harmful outputs. In this paper, we study the bias in LLMs across different dimensions, including gender, race, and religion. We compare the bias in different models and find that while some models are more biased than others, all models exhibit some level of bias. We also find that the bias is often subtle and hard to recognize, making it difficult to detect and mitigate. Our findings suggest that while LLMs can be useful tools, they must be used with caution and awareness of their limitations.

- [Surface Fairness, Deep Bias: A Comparative Study of Bias in Language Models](#), June, 2025
- Looks fair on the surface, but isn't — AI can give answers that seem balanced, yet still hide subtle bias beneath.
- Grades people differently — AI has shown it can rate one group's answers higher than another's, even when the facts are wrong.
- Gives unequal advice — AI can suggest better deals or terms for some groups over others, even without being told who they are.
- Remembers and repeats bias — Once AI "learns" about someone, it can carry that bias into future answers without you realizing it.

25

Strategies for Avoiding Bias in Prompts

Check Your Prompt First • Bias often enters before the AI responds • Review your own wording for assumptions • Use factual inputs, not interpretations	Use Neutral Framing • Use "Party A / Party B" instead of labels that imply roles or status • Frame issues as goals or needs, not stereotypes • Language descriptive not evaluative
Avoid Assumptions • Focus only on information actually provided • Don't embed motivations, emotions, or intent • Center prompts on goals, interests, and facts	Tell the AI to Be Inclusive • Add explicit instructions to avoid bias • Example: "Create an equitable, inclusive summary that avoids assumptions and reflects all parties' perspectives fairly." • Use this in summaries, issue lists, and brainstorming prompts,

26

EXAMPLES: Neutralizing Prompts to Avoid Bias

"Summarize why Party B's delay created the current dispute and outline the problems it caused for Party A."

"Create a neutral summary of the dispute, outlining each party's description of the delays and their impact, based solely on the information provided and without assigning responsibility."

"Break down the strengths of Party A's position so I can prepare for mediation."

"Provide a neutral analysis of the key issues raised by all parties, based solely on the information provided, without evaluating the merits or offering legal conclusions."

"Summarize how the employee misinterpreted the company's policy and why that led to the conflict."

"Summarize the differing interpretations of the policy as described by each party, outlining the issues that contributed to the conflict without indicating which interpretation is correct."

27

Strategies for Checking Bias in Outputs

Ask the AI to reflect on its assumptions:

"What assumptions did you make in this response about gender, roles, or priorities?"

Iterate to request a bias-aware revision:

"Revise this response to remove any unintended bias and reflect inclusive, neutral language."

28

THE PROBLEM

Security - Confidentiality & Privacy



+



≠



29



AI Platform Comparison for Confidentiality & Pricing (2025)

Plan Type	ChatGPT	Gemini	Microsoft Copilot
Consumer	Free or <u>Plus \$20/mo</u> ✗ Not safe for client	Free or Advanced <u>\$20/mo</u> ✗ Not safe for client	Bing Copilot Free ✗ Not safe for client
Business/Team	Team <u>\$25–30/user/mo</u> ✓ Safer, no training	Workspace <u>\$20–30/user/mo</u> ✓ Safer, no training	M365 Copilot <u>\$30/user/mo</u> ✓ Data stays in tenant
Enterprise	Custom <u>\$60+/user/mo</u> ✓✓ SOC 2, contracts	Enterprise <u>\$30+/user/mo</u> ✓✓ Compliance-grade	M365 E3/E5 + <u>\$30/user/mo</u> ✓✓ SOC 2, HIPAA-ready

Legend: ✗ Not safe for confidentiality | ✓ Safer (with limits) | ✓✓ Safest, enterprise-grade

30

SOLUTION

Security - Confidentiality & Privacy



legal.Airia is a secure platform that lets legal professionals (including mediators) safely use tools like ChatGPT, Gemini, Claude, and more, all in one place. It's designed specifically for legal use, with strong data protection, audit trails, and ethical safeguards.

31

Dealing with Confidentiality and Privacy When Using AI Tools

1. NO PII

Use "Party A/B" and general terms — never names, addresses, or other identifiers.

ALWAYS ANONYMIZE**3. No Live Case Data**

Never input real-time case details unless on a secure, enterprise platform.

2. Turn Off Memory

Disable ChatGPT memory or use incognito mode for sensitive content.

4. Structure, Don't Store

Use AI to design processes and explore ideas — not to retain confidential info.

32



IN THE ROOM DEMOS AND EXERCISES

33

WHAT THIS MIGHT LOOK LIKE:

1. Option Generation
2. Risk Analysis Support
3. Visual Aids + Reality-Testing Tools
4. Refined Option Generation
