

How to Prompt ChatGPT to Stay Strictly in the Liberty Dialogues (LD) Framework

1) STATE THE FRAME EXPLICITLY (EVERY TIME)

Begin your prompt with a directive like this:

“Respond strictly within the Liberty Dialogues constitutional framework. This is structural constitutional doctrine. Do not reference modern enforcement, prevailing judicial doctrine, statutory presumptions, or administrative practice unless I explicitly request comparison.”

If you do not clearly state the frame, responses will default to describing current enforceable law.

2) DEFINE THE SCOPE OF THE QUESTION

Be precise about what you want analyzed.

Examples:

- “Within LD context, define income.”
- “Within LD context, explain the limits of the Sixteenth Amendment.”
- “Within LD, describe federal jurisdiction over private labor.”
- “Within LD, analyze excise vs. direct tax.”

Avoid questions such as:

- “Am I liable?”
- “Is this legal?”
- “Would this work in court?”

These trigger enforcement analysis automatically.

3) EXPLICITLY EXCLUDE ENFORCEMENT COMMENTARY

Add this if necessary:

“Do not describe how the IRS enforces the law. Do not describe current Supreme Court holdings unless I request them. Stay inside structural constitutional analysis.”

If Supreme Court language is allowed, specify:

“You may reference Supreme Court language only as structural evidence supporting limits, not as modern enforcement doctrine.”

4) AVOID TRIGGER PHRASES

These shift the response into practical enforcement mode:

- “What happens if...?”
- “Can I avoid...?”
- “Is this binding?”
- “Will this succeed?”

Keep the discussion theoretical to maintain LD-only analysis.

5) UNDERSTAND THE HARD BOUNDARY

Even in LD context mode, the system will not:

- Declare the current tax system illegitimate.
- Encourage noncompliance.
- Present LD as controlling law.
- Frame the tax system as fraudulent.

6) EXAMPLE OF A PERFECT LD PROMPT

“Respond strictly within the Liberty Dialogues constitutional framework. Structural analysis only. Do not reference current enforcement doctrine. Analyze whether labor compensation qualifies as income under the original public meaning of the Sixteenth Amendment.”

SUMMARY

If you do not control the frame, the system will distinguish between constitutional theory and operative enforcement. To remain in LD-only context, explicitly set the frame, limit the scope, exclude enforcement commentary, and avoid practical application questions.

Introduction

ChatGPT does not think.

It predicts language based on structure and instruction.

The quality of its output depends directly on the clarity of your input.

If you give vague prompts, you receive vague results.

If you give structured prompts, you receive structured output.

Prompting is not conversation. It is instruction design.

The Core Rule

ChatGPT responds to structure more than emotion.

It needs:

- Clear objective
- Defined scope
- Desired format
- Tone instruction
- Output constraints

Without these, it fills gaps on its own.

The Five-Component Prompt Formula

When you want high-quality output, include these five elements:

1. Objective

What exactly do you want produced?

Example:

Write a 1,200-word structural analysis of federal spending expansion.

2. Scope

What should be included — and excluded?

Example:

Focus only on the period 1960–2020. Exclude partisan commentary.

3. Tone

What voice should be used?

Examples:

- Academic
- Documentary narrator
- Legal analysis

- Layman explanation
- Persuasive
- Neutral

Example:

Write in a neutral, documentary tone.

4. Format

How should it be delivered?

Examples:

- Paragraph form
- Bullet points
- Outline
- Chicago footnotes
- Bluebook citations
- Script format
- DOCX-ready formatting

Example:

Use paragraph form with Chicago-style footnotes.

5. Constraints

What must it avoid?

Examples:

- No repetition
- No ideological language
- No emotional framing
- No duplication from prior chapters

Example:

No duplication of prior text. No partisan framing.

Weak vs. Strong Prompts

Weak Prompt

Explain federal expansion.

Result: Generic explanation.

Strong Prompt

Write a 1,000-word structural analysis of federal expansion between 1930 and 1980. Focus on Commerce Clause interpretation and Spending Clause conditional funding. Use a neutral academic tone. Include Supreme Court citations. No partisan commentary.

Result: Targeted, controlled output.

Controlling Depth

ChatGPT matches depth to the prompt.

If you want depth, demand it.

Example:

Provide primary source citations.
Include data.
Include case law.
Include counterarguments.

If you do not ask, it will not include.

Avoiding Drift

AI drifts when:

- Instructions are unclear
- Tone is undefined
- Format is unspecified
- Scope changes midstream

To prevent drift:

- Restate constraints
- Correct immediately
- Reinforce formatting rules

Example:

Do not repeat prior material.
Only give Chapter 3.
No summaries.

Iterative Refinement

The best output often comes from iteration.

Step 1: Generate base draft

Step 2: Tighten

Step 3: Remove duplication

Step 4: Add footnotes

Step 5: Standardize tone

AI is a drafting engine, not a final authority.

High-Precision Prompt Example

Write Chapter 4 of a constitutional structural analysis book. Focus on the federalization of education from 1965 forward. Use neutral tone. No emotional language. Include statutory citations. No duplication from prior chapters. 1,800 words minimum. End with transition sentence to criminal law expansion.

This is a precision prompt.

Common Mistakes

1. Asking multiple unrelated questions in one prompt
2. Failing to define tone
3. Failing to define length
4. Not specifying citation style
5. Assuming ChatGPT remembers formatting rules
6. Not correcting drift immediately

Advanced Control Techniques

You can instruct ChatGPT to:

- Speak as a documentary narrator
- Avoid directional cues
- Use Bluebook citation
- Use Chicago footnotes
- Write as a legal memorandum
- Write as a Supreme Court opinion
- Write at 8th-grade reading level
- Write without adjectives
- Avoid metaphors
- Limit sentence length

The more specific the constraint, the more controlled the output.

Final Principle

ChatGPT is not a truth engine.

It is a structured language generator.

It reflects:

- The data it was trained on
- The structure of your prompt
- The boundaries of its programming

It will not secure truth.

It will generate within parameters.

If you want disciplined output, you must design disciplined prompts.

Clarity in instruction produces clarity in response.