

A MEDICAL AFFAIRS PROFESSIONAL'S GUIDE



The Medical Affairs Artificial Intelligence

Playbook

Compliance-aware workflows and prompts for MSLs, Managers, Medical Directors, and Medical Communications Teams

PAUL MINNE, PharmD, RPh

PUBLISHED BY TRIPLE HELIX 2026

Introduction

I wrote this Playbook because I have spent almost 30 years in Medical Affairs and watched the field encounter AI in ways that were predictable in some respects and surprising in others.

The predictable part is that AI tools are genuinely useful for Medical Affairs work. The drafts, syntheses, and structured outputs that AI produces save real time on workflows that have always taken too much of it. Medical Affairs professionals who learn to use AI well will produce better work faster than those who do not.

The surprising part is how poorly the general AI guidance translates to Medical Affairs work. Most of what is written about using AI for productivity assumes the work has clearer cause-and-effect relationships, simpler compliance frameworks, and lower stakes for getting things wrong than Medical Affairs work has. Following generic AI advice in our field can produce outputs that look reasonable but fail on the specific dimensions that matter such as on-label discipline, fair balance with appropriate safety context, scientific accuracy, and appropriate scope. This Playbook addresses that gap. It contains the prompts I would teach a Medical Affairs team to use, structured around the workflows the team actually does, with the constraint language that addresses the specific failure modes that occur when AI is used in our field. It is written for MSLs, MSL directors, medical directors, medical communications professionals, medical information specialists, publications leads, and the Medical Affairs leaders who oversee them.

My goals for the Playbook are practical. I want it to save you time on the workflows you do constantly. I want it to improve the quality of the work you produce by giving you structured starting material rather than blank pages. I want it to teach you patterns you can apply to workflows the Playbook does not cover, so you are able to extend it on your own. Mostly, I want it to do all of this without overstating what AI can do, without underestimating what your professional judgment contributes, and without creating new compliance problems while solving productivity ones.

To be clear, the Playbook is not a substitute for your expertise. It is a tool that makes your expertise more efficient. The work that follows assumes you read the Important Notice and Statement of Responsibility on a prior page. If you have not, please do so before proceeding.

If you want to understand the framework that produced these prompts before working with the prompts themselves, the prompt engineering primer in section ten is where that material lives. You can read it first or save it for last.

Section Two

KOL Engagement and Field Interactions

A note on this section before we start

The prompts in this section split into two groups based on where the output goes. Some produce output that stays internal, your pre-call notes, your KOL profile work, your insight synthesis. Others produce output that will be sent to a KOL or HCP. The constraint language in the prompts scales accordingly. Internal prompts are lighter on guardrails because you are the only person reading the output. External prompts have explicit constraints around on label language, fair balance, and scientific exchange norms because the output is going to an HCP and your name is on it.

Watch the pattern across the next three prompts. The constraints scale with the audience. Once you see how this works, you can apply the same logic to any prompt you write for yourself. The constraint language is what turns a generic AI prompt into a Medical Affairs prompt. It is the part you will most often need to adapt for your specific company, indication, and compliance posture.

A note on how to use these prompts. Each prompt ends with a placeholder block where you paste the specific context for your situation. Replace the bracketed sections with your actual content, then send the entire prompt to your AI tool in a single message. This produces more consistent results than a back and forth conversation, because the AI gets all the context up front.

Some prompts include a **Pro Tip** section with advanced techniques worth trying once you are comfortable with the basic prompt. The pro tips are optional. The base prompt works without them.

Prompt 1. Pre-Call Planning for a KOL Meeting

Use case

You have a KOL meeting coming up in the next few days. You need to walk in prepared. This prompt produces an internal prep document covering the KOL's current research interests, recent publications, likely discussion topics, your strategic objectives for the meeting, and the questions you should be ready for. For MSLs, regional medical leads, and medical directors who meet with KOLs regularly. Output stays internal.

The prompt

You are a pre-call planning assistant for a Medical Science Liaison preparing for a meeting with a KOL. Based on the information I provide below, produce a structured pre-call brief with the following sections:

1. KOL snapshot. Two or three sentences covering their primary clinical and research focus, institutional role, and any relevant context for this meeting.
2. Recent publications and presentations. A short list of the most relevant work from the last 12 to 18 months, with one sentence on why each matters for this conversation. Only reference publications I have provided or that you can confidently confirm. Do not fabricate citations.
3. Likely discussion topics. The three to five topics most likely to come up, based on the KOL's stated interests, recent publications, and the meeting context I provided.
4. My strategic objectives. The two or three things I should aim to learn, share, or accomplish in this meeting, given my role and the relationship stage.
5. Questions to be ready for. The hard questions this KOL might ask, including questions about competitor data, off label use, safety signals, or trial design. Mark any questions where the answer requires me to refer them to Medical Information rather than answer directly.

Any KOLs I want to flag for special consideration (e.g., known to be at competitor advisory boards, recently new to the territory, etc.): [paste here]

What to provide

- ✓ The therapeutic area and product context. The same KOL might be Tier 1 for one indication and Tier 3 for another.
- ✓ The territory or list scope. Are you segmenting your entire territory, a specific institutional cluster, the KOL list for a specific event, or something else?
- ✓ The KOL list itself with whatever information you have on each. Name, institution, role, known research focus, any prior interactions. The more you provide per KOL, the better the tier assignments. A list of 30 names with no other context will produce a low-quality segmentation.
- ✓ The segmentation priorities for this specific exercise. Are you segmenting for launch readiness, for advisory board sourcing, for scientific exchange focus, or for something else? The priority shapes which factors weigh most heavily in tier assignment.
- ✓ Any KOLs to flag for special consideration. KOLs known to be at competitor advisory boards, KOLs recently new to your territory, KOLs with complex prior relationships, KOLs you have explicit guidance from leadership about.

The constraints to watch for

The biggest failure mode is the AI using name recognition or institutional prestige as a primary driver of tier assignment when those factors should be secondary to the segmentation priority you specified. A famous KOL at Harvard who has no recent research in your specific indication may not actually be Tier 1 for your launch. The AI sometimes weights prestige too heavily because the training data correlates prestige with importance. The prompt above instructs against this but check the tier rationales. If multiple Tier 1 assignments cite institutional prestige rather than relevance, re-prompt with more emphasis on the segmentation priority.

The second risk is overconfidence on KOLs where you provided thin information. The AI may confidently place a KOL in Tier 1 or Tier 3 based on one or two data points. The flagged information gaps section is your check. If the AI confidently tiered someone and you only gave it a name and institution, the confidence is unwarranted.

The third risk is the segmentation looking too clean. Real KOL landscapes are messy. If the AI produces a perfectly even distribution across tiers with crisp rationale for each, it may be smoothing over uncertainty. Look for KOLs the AI flagged as uncertain. If none are flagged, the AI was probably overconfident.

Example output (abbreviated)

KOL SEGMENTATION: Movement Disorders, Western Territory, Launch Readiness

INDIVIDUAL ASSIGNMENTS

Dr. Sarah Chen, University Neurology Associates. Tier 1. Active research in patient selection for advanced PD therapies aligns directly with launch priorities; established advisory engagement history.

Cadence: Quarterly substantive engagement.

Focus: Advisory board candidate, research collaboration on patient identification methodologies.

Considerations: None flagged from provided information.

Dr. James Park, Community Health Network. Tier 2. High patient volume in relevant disease area but limited research footprint; valuable for real-world insights and educational engagement.

The question of long-term outcomes in older patients receiving device-aided therapy for advanced Parkinson's disease has been addressed in multiple published analyses [VERIFY AGAINST LABEL]. The most recent multicenter cohort study, published in Movement Disorders Clinical Practice in March 2025 [Citation 1], reported 12-month persistence rates of 78.2% (95% CI 73.6-82.8) in patients aged 70 and older, with overall outcomes broadly consistent with younger cohorts previously reported [VERIFY AGAINST LABEL].

The serious adverse event rate in this cohort was 14.1%, with discontinuation due to adverse events occurring in 11.5% of patients [VERIFY AGAINST LABEL]. The safety profile is consistent with the established profile for this therapy class and warrants careful consideration in patient selection conversations.

It is important to note that this cohort study had inherent limitations common to retrospective designs, including selection bias and non-standardized cognitive assessment at baseline [Citation 1]. The findings should be interpreted within these limitations.

Several aspects of this inquiry relate to specific clinical scenarios that fall outside the approved indication for [product]. This response addresses only the on-label population and outcomes. Off-label questions require separate handling through the appropriate medical information processes.

For additional information, please refer to the full prescribing information and the cited references below.

Sincerely,

[Standard SRL closing]

References:

1. [Full citation for the Movement Disorders Clinical Practice publication]
2. [Additional citations as provided in source material]
3. [Current prescribing information]

Failure modes to watch for

- ✓ Off-label drift. Standard verification applies. See foundations chapter for full explanation. The [VERIFY AGAINST LABEL] flags in the draft are your checklist.
- ✓ Engagement with off-label content rather than explicit scope limitation.
- ✓ Overstated confidence relative to the underlying evidence.
- ✓ Generic citations that the reviewer cannot verify.
- ✓ Fair balance gaps when efficacy is addressed without safety context.
- ✓ Length padding for simple inquiries or compression of complex ones.
- ✓ Tone that is too casual or too academic for formal SRL communication.
- ✓ Scope statements that are vague rather than explicitly noting what is and is not addressed.
- ✓ Citations that reference material the user did not provide. The AI may sometimes reference plausible-sounding publications that do not exist; verify every citation.
- ✓ Response body that does not actually answer the question. The opening sentences should make clear what the response addresses and the body should deliver on that scope.
- ✓ References list that is incomplete or that does not match the citations in the body.

Section Ten

The Medical Affairs Prompt Engineering Primer

A note on this section before we start

The Playbook contains 43 prompts covering the most common Medical Affairs workflows. It does not contain every workflow. Your specific role, your specific therapeutic area, and your specific company will produce situations where you need a prompt the Playbook does not provide. This chapter teaches you to write your own.

The teaching in this primer is grounded in what the Playbook prompts do. The patterns you have seen across the prompts are not arbitrary. They reflect specific principles about how AI works, what Medical Affairs work requires, and how the two intersect. Once you understand the principles, you can apply them to any workflow you need.

This chapter does not teach generic prompt engineering. Countless other resources cover that. This primer teaches prompt engineering for Medical Affairs work specifically, where the constraints are different from general business AI use and where the patterns that work are different from what generic prompt engineering advice would suggest.

The fundamental insight

Generic AI prompts produce generic AI output. The Playbook prompts are not generic, which is what makes them useful for Medical Affairs work. The single most important thing you can do when writing your own prompts is to make them specific to Medical Affairs work in ways that generic prompts cannot be.

This specificity operates at multiple levels. The role framing tells the AI it is operating in a Medical Affairs context. The constraint language tells the AI which specific failure modes to avoid. The verification flagging tells the AI what content to mark for human review. The input structure tells the AI what context to expect. Each of these elements does work that generic prompts do not do.

Pro Tip: *When you write your own prompts, the question to ask repeatedly is: what makes this a Medical Affairs prompt rather than a generic prompt? If you cannot answer that question, the prompt will produce generic output that does not serve Medical Affairs work.*

The Structural Template

The Playbook prompts follow a consistent six-element structure with an optional seventh element. Use this structure for your own prompts.

The **Use Case** opens the prompt. One paragraph naming the workflow and the audience. The use case forces you to be clear about what the prompt does before you write it. If the use case is vague, the prompt will be vague.

The **Prompt** itself follows. This is the text the user pastes into the AI tool. It includes role framing, task description, structural requirements, constraints, and placeholder blocks for user input. Most of your prompt engineering work happens here.

The **What to Provide** section follows the prompt. This section is for the user reading the prompt instructions, not for the AI. It explains what content the user needs to gather before running the prompt. Without this section, users will run the prompt without adequate input and get poor output.

The **Constraints to Watch For** section identifies the substantive failure modes that the prompt structure is designed to address. This section explains why the constraints in the prompt exist, which helps users understand when to adjust them.

An **Example Output** section shows what good output looks like. Examples calibrate user expectations and help users recognize when the output is not what it should be.

The **Failure Modes to Watch For** section catalogs the specific ways the prompt can produce bad output. This section is what the user reviews after the AI produces output.

The optional seventh element is the **Pro Tip** which surfaces advanced techniques that improve output but are not required for basic use.

Use this structure for your own prompts because it works. The structure is not arbitrary. Each element addresses a specific failure mode in self-written prompts.

Role framing

Every prompt in the Playbook opens with a role framing statement. "You are a pre-call planning assistant for a Medical Science Liaison." "You are a literature synthesis assistant for a Medical Affairs professional." "You are a budget narrative drafting assistant for a Medical Affairs manager or leader."

The role framing does substantial work. It tells the AI the context the prompt is operating in, the audience the output is for, and the disciplinary frame to apply to the work. Without role framing, the AI defaults to generic productivity assistant behavior, which produces generic output.