



How To Properly Speak To AI

Table of Contents

Introduction.....	1
Part 1: The core principles	2
Part 2: Techniques that reliably improve output	2
Part 3: Giving AI good inputs, not just good words	3
Part 4: Common mistakes worth avoiding.....	4
Part 5: A quick before-and-after	4
Quick-reference checklist	5
Closing thought	5
Let us find the right AI for you.....	5

Introduction

Modern AI assistants are capable generalists, but the quality of what they give back is still shaped heavily by the quality of what you put in — this is often called “prompt engineering,” though in practice it’s less like engineering and more like learning to give clear instructions to a very capable, very literal new hire. The core idea driving nearly all current guidance, from Anthropic, OpenAI, and independent researchers alike, is the same: the model isn’t reading your mind, it’s reading your words, so vague input reliably produces vague or generic output. This guide lays out what currently works, what to avoid, and how to think about the inputs you give AI beyond just the words you type.

Part 1: The core principles

Nearly every credible guide to prompting converges on the same handful of fundamentals.

Be specific and explicit. Don't assume the model will infer what you want — state it directly. “Write a marketing email” produces something generic; “write a 150-word email announcing a 20%-off weekend sale to existing customers, in a warm but not pushy tone” gives the model enough to work with on the first try.

Give context, not just an instruction. The single most common mistake identified across current guidance is assuming the AI already has context it doesn't have. Telling an assistant to “summarize the meeting” without saying what matters in that summary, who it's for, or what decisions were made leaves it guessing — and it will often guess wrong. A sentence or two of background is rarely wasted effort.

One task at a time. Asking a single prompt to summarize an article, suggest improvements, write a social post about it, and generate an image idea all at once splits the model's attention across competing objectives, and quality drops across all of them. Breaking a request into sequential, focused prompts consistently produces better results than one overloaded one.

State the output format you want. Ask directly for a table, a numbered list, a specific word count, or a particular structure. The model won't guess your preferred format — say it up front rather than reformatting after the fact.

Treat the first response as a draft to refine, not a final answer. Prompting is inherently iterative: give an initial instruction, look at what comes back, and adjust — tightening the tone, adding a missing constraint, or correcting a misunderstanding — rather than starting over from scratch. This iterative loop is consistently identified as one of the highest-leverage habits in current guidance.

Part 2: Techniques that reliably improve output

Beyond the fundamentals, a smaller set of specific techniques shows up across current guidance from Anthropic, OpenAI, and independent research.

Show examples instead of over-explaining. If you want a particular style, structure, or tone, pasting in one or two examples of what “good” looks like is often more effective than a long paragraph of instructions trying to describe it in the abstract. This is sometimes called “few-shot” prompting. That said, more isn't always better — most current guidance suggests two to four well-chosen examples for a given task, since piling on too many can cause the model to overfit to surface details of the examples rather than generalizing the underlying pattern.

Assign a role when it's genuinely useful. “Act as a patient tutor” or “act as a skeptical editor reviewing this for weaknesses” can usefully shift tone and framing for open-ended, creative, or advisory tasks. It has little effect on tasks with a single correct answer — a factual lookup or a data-classification task won't improve just because you assign the AI a persona, so save this technique for the tasks where framing genuinely matters.

Frame instructions positively rather than negatively. “Only reference numbers that appear in the source document” tends to outperform “don't make up numbers that aren't in the source.” Telling a model what not to do still requires it to process the concept you're trying to avoid; stating the desired behavior directly is generally more reliable.

Give explicit permission to say “I don't know.” Models are more likely to guess confidently than to admit uncertainty unless you tell them that's acceptable. Adding a line like “if you're not sure, say so rather than guessing” measurably reduces confidently-wrong answers on ambiguous questions.

Use step-by-step reasoning selectively, not by default. Asking a model to “think step by step” or “show your reasoning before answering” can improve accuracy on multi-step problems for standard chat models. However, current-generation reasoning models (the models with built-in extended thinking, sometimes labeled “reasoning” or “thinking” modes) already do this internally — explicitly instructing them to think step by step on top of that adds little benefit and can slow the response down for no real gain. Save explicit step-by-step instructions for standard models and simpler tasks.

Use structure for complex prompts. For longer or more complex requests — multiple pieces of context, several instructions, or reference material — separating the parts clearly (labeled sections, a simple numbered list, or tags like `Context:`, `Task:`, and `Format:`) helps the model parse what's an instruction versus what's background material, and reduces the chance it conflates the two.

Part 3: Giving AI good inputs, not just good words

Prompting well is about more than phrasing — what you feed the model matters just as much.

Paste in the actual source material. If you want a summary, a rewrite, or an analysis, paste in the real document, email, or data rather than describing it from memory. A model working from your imperfect recollection of a document will produce a less accurate result than one working from the document itself.

Upload files directly when you can. Most current chat assistants can read an uploaded document, spreadsheet, image, or screenshot directly, which is usually more reliable than manually retyping or summarizing that content into the chat yourself — retyping introduces exactly the kind of information loss you're trying to avoid.

Keep irrelevant material out. Pasting in an entire lengthy document when only one section is relevant can dilute the model's attention, particularly with information buried in

the middle of a long input — a well-documented weak spot for large language models. Trim inputs to what’s actually relevant to the task when you can.

Be precise with data and numbers. If a task depends on exact figures, dates, or names, provide them explicitly rather than expecting the model to recall or infer them correctly — this is one of the more common sources of subtly wrong output.

Keep a working prompt once you’ve refined it. If a prompt for a recurring task (a weekly report format, a standard email reply, a specific analysis) took a few rounds of iteration to get right, save it rather than reconstructing it from memory each time. Treating a reusable prompt as something worth keeping and refining, rather than a one-off, is one of the more practical habits separating casual use from consistently good results.

Part 4: Common mistakes worth avoiding

1. **Assuming shared context that doesn’t exist.** The model only knows what’s in the conversation (and, if it’s searching, the open web) — not your internal shorthand, your company’s specific processes, or a decision made in a meeting it wasn’t part of.
 2. **Bundling too many objectives into one prompt.** Multiple unrelated asks in a single message tend to produce shallower results on all of them rather than a strong result on each.
 3. **Vague quality requests.** Asking for “a professional-sounding” post or email without an example or a specific description of what that means tends to produce generic, forgettable output.
 4. **Piling on excessive examples.** Beyond a handful of well-chosen examples, additional ones tend to add noise rather than clarity.
 5. **Adding “think step by step” out of habit on every model.** It helps on some tasks and models, and is dead weight — or mildly counterproductive — on current reasoning-focused models that already reason internally.
 6. **Treating the first answer as final.** Skipping the refinement step and either accepting a mediocre first answer or abandoning the tool entirely both throw away most of the value a quick follow-up correction would have delivered.
-

Part 5: A quick before-and-after

Before: “Write something about our new product.”

After: “Write a 200-word product announcement for [PRODUCT NAME], a [WHAT IT DOES] aimed at [AUDIENCE]. Highlight [KEY BENEFIT] and [SECOND BENEFIT]. Tone: confident but not salesy. End with a one-line call to action. Here’s an example of the tone we use in past announcements: [PASTE EXAMPLE].”

The difference isn’t length for its own sake — it’s that the second version removes nearly every ambiguity the model would otherwise have had to guess at: topic, audience, length, tone, structure, and a concrete example of “good.”

Quick-reference checklist

Before you send a prompt, check that you've...
Stated the specific task, not just a general topic
Given relevant background the model wouldn't otherwise know
Kept the prompt focused on one main objective
Said what format or length you want the answer in
Pasted in real source material rather than describing it from memory
Included an example, if tone or style genuinely matters
Planned to review and refine the first response rather than treating it as final

Closing thought

None of this requires a technical background — it's closer to the difference between giving a new colleague a one-line task and giving them a clear, well-scoped brief. The habits that separate a frustrating AI interaction from a genuinely useful one are almost all about clarity: being specific, providing real context and source material, asking for one thing at a time, and treating the first answer as a starting point rather than a verdict.

Let us find the right AI for you!

This guide should help you get noticeably better results out of any AI assistant you use, but prompting well is a skill that keeps paying off as you apply it to more of your work. If you are ready to go further and bring well-designed AI workflows into your business, we are here to help.

Contact us today through email or on our website to schedule your free 30-Minute Introduction Call to see what AI can do for you!

Email: Sales@IntelliTechDynamics.com

Website: www.IntelliTechDynamics.com

