

Knowledge Engineering: Six Critical Questions for Knowledge Production in the Age of Artificial Intelligence

Introduction: Not Prompt Engineering, but Knowledge Engineering

Interactions with large language models frequently begin with a measure of disappointment. A user poses a complex and potentially productive question, only to receive a response that is remarkably bland, hollow, or beside the point. This disparity is particularly acute in the domain of serious knowledge production: researchers who seek to employ artificial intelligence to analyse imperialist policies, document popular struggles, and disseminate counter-hegemonic narratives find that its outputs consistently fall short of expectations. Many conclude from this experience that artificial intelligence is unsuitable for rigorous intellectual work.

This conclusion, however widespread, rests on a misdiagnosis. The problem lies not in the tool, but in how it is used.

‘Prompt engineering’ has become the dominant framework for addressing this challenge, focusing on how to formulate language in order to elicit better responses from large language models. Consider the widely adopted [CO-STAR framework](#) — Context, Objective, Style, Tone, Audience, Response — in which four of six elements address matters of ‘form of expression’: how to say it, in what tone, to whom, and in what format. This is not without utility — form of expression matters, and we shall address it in due course. However, from the standpoint of serious knowledge production, this framework exhibits systematic blind spots: epistemological framework, information reserve, methodology, opinion and insight — dimensions that are decisive for the quality of any knowledge product — remain entirely absent from its purview.

To reduce the challenge of employing artificial intelligence to a set of techniques for ‘how to phrase things’ constitutes a fundamental misdiagnosis. A more accurate assessment is this: the root cause of most failed interactions with artificial intelligence lies not in the wording of the prompt, but in the user’s failure to systematically organise their own questions, contexts, positions, and purposes. In other words, this is not a problem of prompt engineering; it is a problem of knowledge engineering. AI in this framework functions as an analytical instrument that produces text as its output — not a writing assistant that generates analysis as a side effect. As Herbert Simon demonstrated with his concept of [‘bounded rationality’](#), human decision-making is constrained not by a lack of information but by the cognitive limits on processing it — a principle that applies with equal force to human-AI collaboration, where the quality of outcomes depends less on the model’s capabilities than on the researcher’s capacity to organise the problem.

Amara is a researcher at a progressive policy institute in Dar es Salaam, tasked with producing a report on why Tanzania’s natural gas and mining boom has not translated into broad-based development. She has a laptop, decent internet, access

to an AI tool, six weeks, no research assistant, and a small budget. She does not code. When she first asks her AI tool to ‘write a report on Tanzania’s extractive sector’, she receives a generic overview indistinguishable from a consultancy brochure — accurate in its broadest outlines, useless for her purposes. The six questions that follow chart the path from that initial disappointment to a genuinely useful research product.

This article poses six critical questions that constitute a framework for the practice of knowledge engineering: What problem are we actually trying to solve? From what standpoint do we view the world? What does the AI need to know? How should the analysis proceed? What should the final product look like? And what is the researcher’s distinctive contribution? These six questions apply to any mainstream large language model, independent of any specific tool — they concern not technology, but how researchers themselves think about and organise the production of knowledge.

One: What Problem Are We Actually Trying to Solve?

Most disappointing interactions with artificial intelligence share a common characteristic: the user proceeds directly to specifics without establishing a strategic purpose. ‘Help me write an article about the US-China tech war’ — such a request is virtually guaranteed to produce a superficial survey, because the request itself is superficial. Artificial intelligence is not incapable of producing work of analytical depth; rather, it cannot determine in which direction the desired ‘depth’ should extend.

Problem orientation is not ‘providing context’; it is an exercise in intellectual leadership. It requires the researcher, before engaging in any dialogue with artificial intelligence, to answer a fundamental question: What am I actually trying to understand? This question appears simple, yet it is precisely the step that most users omit.

Consider a concrete research project. In a study of the US-China semiconductor technology war, the researchers did not vaguely request an ‘analysis of US-China tech competition’, but instead formulated five progressively deepening core questions: How is the United States employing technology denial regimes to maintain geopolitical hegemony, and with what effects? How is China responding to technological containment, and what does this reveal about its development model? What is the current state of China’s supply chain dominance in critical sectors, and what are its implications? How is the technology war reshaping global supply chains and trade relationships? What are the implications for the Global South?

Beneath each core question lay further sub-questions that demarcated the boundaries and depth of the analysis. Under the question concerning US technology denial, for instance, the research further inquired: What policies and mechanisms are being deployed? How do these measures compare with historical technology containment strategies? What are the immediate and longer-term consequences?

The same artificial intelligence tool, confronted with ‘write an article about the US-China tech war’ on the one hand, and with such a structured five-tier system of questions on the other, produces fundamentally different outputs. The former yields an encyclopaedic listing of information; the latter yields a political economy study with an analytical framework, with stratification, with a point of view. The difference lies not in the tool, but in the question.

Rather than asking the AI to ‘analyse Tanzania’s extractive sector’, Amara formulates core questions: Why has natural gas and mining revenue not translated into industrialisation or improved livelihoods? What fiscal and contractual arrangements govern extractive rents, and whose interests do they serve? How have successive resource nationalism policies — from Magufuli’s mining reforms onward — altered the relationship between foreign capital and domestic accumulation? What role do international financial institutions play in shaping extractive governance? Each question generates sub-questions; together, they transform a vague topic into a stratified research programme.

In more advanced research practice, problem orientation can become still more incisive. In another project in which the author participated, researchers were required, before commencing any investigation, to answer a fundamental question: What is the principal contradiction driving the current state of affairs? This was not a rhetorical gesture but a binding requirement — if you could not identify the principal contradiction, the research would not proceed. This appears to be a demanding standard, yet it compels the researcher, before any dialogue with artificial intelligence, to complete the most difficult and most essential intellectual labour: determining what the research truly seeks to reveal.

The quality of the question determines the ceiling of the answer. A precise, stratified, analytically embedded system of questions is itself a product of intellectual leadership. Artificial intelligence can help explore the various dimensions of a question, can survey relevant literature and data — but determining which questions are worth posing can only come from the researcher's own theoretical formation, political commitment, and insight into the realities under investigation.

Two: From What Standpoint Do We View the World?

In the employment of artificial intelligence for knowledge production, a dangerous misconception is proliferating: many users regard large language models as neutral analytical instruments, expecting them to deliver 'objective' assessments on any given subject. This supposed neutrality, however, is an illusion. The training data for large language models is drawn predominantly from the English-language internet — a domain long dominated by Western ideological perspectives. Research has [demonstrated](#) that even open-source models developed in China exhibit considerable Western ideological bias: geopolitically favouring Western positions, epistemologically marginalising scholars from the South, economically naturalising neoliberal prescriptions, and culturally centring Western historical narratives.

When researchers employ these tools without critical reflection, a subtle yet far-reaching process unfolds:

It is not you who are calibrating the AI, but the AI that is calibrating you.

If you request an analysis of a developing country's economic predicament without specifying an analytical framework, the model will almost inevitably proceed from the assumptions of neoliberal economics — market efficiency, comparative advantage, governance reform — because these constitute the dominant analytical paradigm in its training data. You may believe you have received 'objective analysis'; in reality, you have received the reproduction of the biases embedded in the training data.

The solution is not to attempt to render artificial intelligence 'more neutral' — this is neither possible nor desirable. Every analysis embodies a standpoint; the only question is whether that standpoint has been consciously chosen or unconsciously absorbed. Knowledge engineering requires that, before collaborating with artificial intelligence, you explicitly declare your epistemological framework.

In a research project in which the author participated, researchers were required, before beginning any analytical work, to declare in advance a set of theoretical framework documents as the epistemological foundation of the entire study. These documents included a Marxist political economy framework, an analytical framework based on the [Eight Contradictions of the Imperialist 'Rules-Based Order'](#), a theoretical summary of [hyper-imperialism](#), and a framework for contemporary class analysis. These were not optional references but binding prerequisites: without a declared epistemological framework, the research would not commence.

When Amara asks the AI to analyse why extractive revenues have not produced broad-based development, the model defaults to the 'resource curse' literature — institutional weakness, Dutch disease, governance deficits. The analysis is not wrong in every particular, but it systematically obscures the structural dimensions that matter most: the role of

multinational capital in shaping fiscal terms, the historically conditioned weakness of the Tanzanian state's bargaining position, and the international financial architecture that channels rents outward. She corrects for this by declaring her framework explicitly — providing reference texts on dependency and unequal exchange, on the political economy of extractivism in Africa, and on Tanzanian economic policy since *ujamaa*. The AI's subsequent analysis is not 'more biased'; it is more honest about the standpoint from which it proceeds.

To 'declare an epistemological framework' does not mean simply instructing the artificial intelligence to 'use Marxist analysis' — such directives tend to produce hollow texts saturated with rhetorical platitudes. What proves genuinely valuable is the systematic articulation of the core principles, terminological system, and analytical priorities of the chosen framework. The same project drew a precise distinction regarding how to apply the epistemological framework 'invisibly':

Correct structural analysis proceeds as follows: 'The concentration of advanced chip manufacturing in Taiwan creates leverage that shapes US policy options...' — the epistemological framework guides the analysis towards power structures and material conditions, while the analysis itself unfolds through concrete facts and logical reasoning.

Incorrect rhetorical platitudes proceed as follows: 'The imperialist contradictions of capitalism drive the US ruling class to...' — this is not analysis but the substitution of theoretical terminology for analysis. The framework has been crudely imposed upon the surface rather than deployed as a lens for deeper investigation.

The following is an actual prompt fragment from this project specifying how to apply the epistemological framework 'invisibly':

Apply Ideological Framework Invisibly

The ideological framework should:

- Inform analysis of power, interests, and structures
- Help identify what questions to ask
- Provide conceptual tools for understanding

The ideological framework should NOT:

- Appear as explicit rhetorical terms
- Replace evidence with assumption
- Predetermine conclusions
- Make analysis feel propagandistic

(The above fragment is drawn from the project's analysis methods guide, which runs to approximately 960 words. The total volume of prompts and reference guides for this research project exceeds 27,000 words — these figures, along with the output statistics cited later in this article, are drawn from the internal documentation of research projects conducted by the GSI team. Those interested in obtaining the complete prompts may contact the GSI team.)

This distinction is of paramount importance: the epistemological framework should guide what kinds of questions you pose (Whose interests are being served? How do power relations operate? How do material conditions shape possibilities?), rather than predetermining the answers. The framework instructs the artificial intelligence from what angle to view the world, but the analysis itself must rest upon evidence and logic.

For researchers in the Global South, this question is especially urgent. When Western-centric narratives dominate the training data, and when researchers fail to proactively declare their own epistemological positions — an anti-imperialist analytical lens, a structural attention to relations of dependency, a foregrounding of Southern experience — then every

passage that artificial intelligence produces may unconsciously reproduce the knowledge hegemony of the North. To leave one's epistemological framework undeclared is to unconsciously accept the ideological standpoint embedded in the model.

It must be emphasised that this is not a matter of making artificial intelligence 'biased' — on the contrary, it is a matter of making it honest. Every analysis has a perspective; every study has a standpoint. Declaring an epistemological framework transforms an implicit, unconscious bias into an explicit, conscious choice of position. This constitutes the deepest expression of subjectivity in knowledge production: you must know from what standpoint you view the world before you can require the machine to assist your analysis from that standpoint.

Three: What Does the AI Need to Know?

A prevalent misconception treats large language models as a kind of oracle — you pose a question, and it delivers an answer from an omniscient internal knowledge base. In practice, large language models are reasoning engines that operate upon the information provided to them. When input information is sparse, they can produce only generic, hollow responses; when input information is rich and specific, their reasoning capacities can be brought fully to bear. The quality and scale of the information reserve directly determine the depth of what artificial intelligence can produce.

The information reserve has two complementary dimensions: the researcher's own accumulated knowledge, and new information obtained in real time through the internet.

Domain Knowledge: What You Bring to the AI

Every researcher accumulates, through sustained practice, a substantial body of knowledge assets: theoretical frameworks, core concepts, literature surveys, historical background materials. When collaborating with artificial intelligence, these knowledge assets should not remain idle — they should be furnished to the artificial intelligence as reference documents, forming the foundation for its reasoning and analysis.

A critical principle obtains here: reference materials provided by the researcher must be preserved in their entirety, never summarised or compressed. You may consider a theoretical framework document of several dozen pages excessively long, and wish to have the artificial intelligence 'summarise the key points' before proceeding — this is precisely the wrong approach. You can never know in advance which detail will prove critical during analysis. The context windows of contemporary mainstream large language models are now sufficiently capacious to accommodate large volumes of text; the full exploitation of this capacity is a fundamental practice of knowledge engineering.

Equally important is the separation of instructions from reference materials. Your instructions — what you wish the artificial intelligence to do — should be concise and precise; your reference materials — what you wish it to 'know' — should be comprehensive and complete. Conflating the two either buries instructions beneath reference materials or compresses reference materials unnecessarily.

Information reserves at scale can yield impressive results. In a study on the [*80th Anniversary of the Victory in the World Anti-Fascist War*](#), researchers constructed an information reserve spanning ten thematic sections — from the economic foundations of the belligerent powers to the architecture of post-war betrayal — ultimately producing a 5,900-word analytical study containing 150 citations across military archives, economic histories, and diplomatic records in multiple languages. Output at this depth and density is possible only when supported by a rich information reserve.

Amara lacks a large institutional library, but she has knowledge assets from her training and prior work — and these prove decisive. She uploads the full text of Tanzania's 2009 Mining Act and its 2017 amendments, the Natural Wealth and Resources Acts, relevant sections of the Five-Year Development Plan, the TEITI reconciliation reports, and two political

economy analyses of East African resource governance her institute had previously produced. She resists the temptation to summarise first; she provides them whole. The difference is not incremental but categorical: with these documents, the analysis cites specific fiscal terms, identifies discrepancies between legislated and effective tax rates, and traces the gap between policy intention and developmental outcome.

Internet Search: Opportunities and Pitfalls

One of artificial intelligence's major capabilities is its ability to search the internet in real time, accessing the most current information. This means that research need not be confined to existing knowledge reserves; it can draw upon the latest policy documents, statistical data, research reports, and news coverage as they appear. However, this capability is accompanied by a distinctive challenge: how do you determine whether the information that artificial intelligence retrieves is credible?

In the US-China semiconductor technology war research project, the researchers specified in advance a credibility hierarchy for data sources: official policy documents (high) → peer-reviewed research (high) → research institution reports (medium-high) → industry analysis (medium) → news reporting (requiring cross-verification) → personal blogs and anonymous social media (excluded). This hierarchy was not decorative — it directly determined the weight that artificial intelligence should assign to different sources during analysis. The project simultaneously required that critical data be corroborated by at least two independent sources.

The deeper problem is that of artificial intelligence 'hallucination' — the fabrication of data that appears plausible but does not in fact exist. This is not an occasional malfunction but a systematic weakness of large language models. The following is an actual prompt fragment from the same project, designed to guide the artificial intelligence in identifying its own hallucinations:

Warning Signs of AI Hallucination

Be especially vigilant for:

1. Plausible-sounding URLs that don't exist
2. Fabricated statistics
 - Very specific numbers that can't be found
 - Statistics that "perfectly" support the argument
3. Invented organisation names
 - Similar to real organisations but slightly different
 - E.g., "China Automobile Association" vs "China Association of Automobile Manufacturers"

(The above fragment is drawn from the project's data verification guide, which runs to approximately 1,300 words.)

Sound knowledge engineering practice demands that systematic scepticism be maintained towards all information returned by artificial intelligence — this is itself an epistemological stance.

Information Reserve as Ongoing Practice

It bears emphasis that the information reserve is not a one-time preparatory exercise but a continuous practice that runs throughout the entire research process. Documents curated in the course of one project — theoretical frameworks, data source guidelines, credibility assessment criteria — become reusable knowledge assets. As projects accumulate, information

reserves grow richer, and the depth of support that artificial intelligence can provide increases correspondingly. This constitutes a positive feedback loop: the greater the investment in knowledge engineering, the greater the returns.

Four: How Should the Analysis Proceed?

Many researchers, when employing artificial intelligence, confine themselves to specifying ‘what to do’ — ‘analyse this issue’, ‘write an article about that topic’ — without specifying ‘how to do it’. This is akin to engaging a highly capable research assistant while providing only a vague direction and no methodological guidance whatsoever. The result is predictable: artificial intelligence organises its analysis according to its own defaults, producing text that appears reasonable but lacks methodological self-awareness.

The fourth question of knowledge engineering requires the researcher to become a project manager: specifying not only the objective, but the path towards it.

Two complementary strategies present themselves. The first is to prescribe the process: to furnish the artificial intelligence with a clear, step-by-step workflow. In the US-China semiconductor technology war research project, the analytical method was explicitly prescribed in four stages:

The first stage was evidence mapping — identifying all verified sources, extracting key facts, and annotating each source’s perspective and credibility. The second stage was pattern identification — seeking recurring themes across sources, identifying areas of consensus and disagreement, and noting changes over time. The third stage was critical analysis — interrogating the relationship between stated justifications and actual outcomes, identifying whose interests are served, and examining ‘what is not being said’. The fourth stage was synthesis — developing coherent arguments from the evidence whilst acknowledging complexity and uncertainty.

More granularly, the project required each argument to follow a five-step construction: claim — evidence — reasoning — qualification — significance. The following is an actual prompt fragment specifying this five-step argument construction:

Argument Construction

For each major argument:

1. Claim: State the argument clearly and specifically
2. Evidence: Present the supporting data with citations
3. Reasoning: Explain why the evidence supports the claim
4. Qualification: Note any limitations or counterarguments
5. Significance: Explain why this matters

Quality Checklist for Arguments

- Is the claim specific and falsifiable?
- Is the evidence from verified sources?
- Is the reasoning logical and explicit?
- Are limitations acknowledged?
- Is significance explained?

(The above fragment is likewise drawn from the project’s analysis methods guide.)

No step may be omitted: you must state clearly what the argument is (claim), what supports it (evidence), why the evidence supports the argument (reasoning), under what conditions the argument holds or does not hold (qualification), and why the argument matters (significance).

This methodological prescription may appear mechanical, but it possesses profound epistemological value. The act of prescribing a process compels researchers to think systematically about their own analytical practice — you are forced to answer explicitly: Where should the analysis begin? Through what stages should it proceed? How is argumentative rigour to be ensured? Many researchers, at the moment they are compelled to articulate their methodology, examine for the first time the analytical habits they have long taken for granted.

Amara prescribes a four-stage workflow. First, map the fiscal architecture — legislated rates, effective rates, stabilisation clauses in mining development agreements — drawing on her uploaded documents and verified sources. Second, identify patterns across time: how have terms shifted between the liberalisation of the 1990s, Magufuli's resource nationalism, and the present? Third, critically analyse the gap between revenue captured and developmental expenditure, interrogating whose interests the prevailing arrangements serve. Fourth, synthesise these findings into an argument about structural impediments to broad-based development. Without this sequence, the AI would have produced a thematic survey; with it, the analysis builds cumulatively, each stage drawing on the last.

The second strategy is task decomposition: breaking complex tasks into manageable sub-tasks. A sweeping research question — such as 'analyse the restructuring of global semiconductor supply chains' — if presented to artificial intelligence as a single request, will inevitably yield a superficial overview. If, however, it is decomposed into more specific sub-tasks — first analyse US export control policies, then analyse China's indigenous substitution efforts, next examine the strategic positioning of third countries, and finally synthesise these analyses into an overall assessment — each sub-task receives more thorough treatment.

In more complex research practice, the methodology can extend to sixty or seventy steps. The 'steelmaning and dialectical fortification' phase alone, for instance, comprises over a dozen steps. Among them is a 'triple hostile reader test': requiring that one's own argument be scrutinised from the perspectives of a capitalist critic, a state-socialist critic, and an anarchist critic, respectively — can you rebut the strongest objections? This degree of methodological refinement ensures that the analysis is not merely internally consistent within its own framework, but can withstand rigorous challenge from opposing standpoints.

An important inflection point is implicit here. From four steps to seventy, the growth in methodological complexity gives rise to a practical problem: a single prompt can no longer accommodate such extensive methodological prescriptions. When researchers must simultaneously manage dozens of analytical steps, each with its own input-output standards and quality requirements, the natural next step is to distribute the methodology across multiple specialised agents — this is precisely the problem that declarative multi-agent systems, which we have discussed in detail in a separate article, are designed to address. Yet even before entering the domain of multi-agent systems, the researcher who prescribes a clear methodology within a single dialogue can already achieve significant improvements in the analytical quality of AI-generated output.

Five: What Should the Final Product Look Like?

The same body of knowledge, expressed in different forms and serving different purposes and audiences, becomes an entirely different thing. A piece of in-depth research analysis, written as an academic paper, takes one shape; rewritten as training materials for community organisers, it takes another; produced as a podcast script, it takes yet another. Traditionally, effective writing has required the researcher to master three elements simultaneously: a solid command of background knowledge, depth of analytical insight, and command of the appropriate register. All three are indispensable,

which explains why capable writers have always been scarce — those who simultaneously excel in analysis and expression are not readily found.

Artificial intelligence alters this equation. If you have addressed the preceding four questions — problem orientation, epistemological framework, information reserve, and methodology — artificial intelligence can provide powerful assistance with form of expression. You need not be simultaneously an analytical expert and a master of prose. This does not mean, however, that you may neglect the specification of form — on the contrary, the more precisely you define your formal requirements, the more closely the AI's output will approximate your expectations.

In the US-China semiconductor technology war research project, the output template prescribed the Tricontinental dossier style with sentence-level precision. The document type was specified as 'Tricontinental-style research dossier, 5,000 to 10,000 words'. Sentence complexity was prescribed through positive and negative examples:

Appropriate sentence complexity proceeds as follows: 'While the United States has framed its semiconductor export controls in terms of national security, citing the potential military applications of advanced chips, the broader economic and geopolitical motivations — maintaining technological supremacy and slowing a rising competitor — are evident in the scope and evolution of these restrictions.'

Excessive simplicity proceeds as follows: 'The US says export controls are about security. But they're really about competition.'

These two sentences convey approximately the same information, yet the former employs complex syntactic structure to express the nuanced relationship between surface justification and underlying motivation, while the latter reduces the complexities of geopolitics to a simple binary opposition. For an academic research report, only the former mode of expression can bear the precision demanded by the analysis.

Vocabulary was likewise subject to detailed specification: terms of analytical precision such as 'hegemony', 'containment', 'industrial policy', and 'technological sovereignty' were to be employed; rhetorical ideological terms were to be avoided — structural analysis was to speak for itself, rather than being replaced by labels.

Amara's institute has asked for a policy report of approximately 8,000 words, aimed at Tanzanian parliamentarians and East African civil society organisations — an audience that demands analytical rigour but not academic jargon. She specifies accordingly: the register should be formal but accessible, employing terms such as 'fiscal leakage', 'beneficiation', and 'local content requirements' where analytically necessary, but avoiding specialised academic vocabulary that would alienate non-specialist readers. She provides a previous institute publication as a style reference and specifies that every factual claim must cite a source. The report should include a summary of recommendations — not because the AI will generate them (that is her work), but because the structural template must accommodate them.

In more complex research projects, researchers prepare multiple writing reference documents specifying, respectively, content structure, citation standards, technical writing conventions, and referencing formats. For example, it may be required not only that sources be cited, but that each endnote contain context and a verifiable claim — enabling readers to independently verify every factual assertion.

It should be noted that the examples above pertain exclusively to textual output. Were the product a video script, a podcast outline, an infographic, or another multimedia form, the corresponding format requirements would naturally differ entirely — yet the same level of detailed specification would be required. The critical point is not what is specified, but that specification is **necessary**. Vague formal requirements produce vague results, just as vague questions produce vague answers.

This question also yields a cumulative benefit. With a solid information reserve and methodology in place, the researcher can rapidly re-express the same body of knowledge in multiple forms. A researcher's deep analysis of a given topic can be disseminated simultaneously through academic papers, policy briefs, training materials, and social media content. Artificial intelligence renders this 'analyse once, express many times' approach a practical reality — and this is of particular significance for resource-constrained research institutions in the Global South.

Six: What Is the Researcher's Distinctive Contribution?

The preceding five questions define the supporting structure of knowledge production — problem orientation prescribes the direction, the epistemological framework prescribes the perspective, the information reserve provides the raw materials, methodology prescribes the process of elaboration, and form of expression prescribes the shape of the final product. Yet if only these five questions were addressed, the entire system would remain inert. It requires an initiator to conceive it, a decision-maker to guide it, and a judge to assess the value of its outputs. This sixth question — what is the researcher's distinctive contribution — is not a step in the knowledge production process but the force that drives and evaluates the entire process.

The cybernetics pioneer Norbert Wiener, in his 1950 work *The Human Use of Human Beings*, offered a profound observation: the first industrial revolution automated physical labour, yet it did not render human beings superfluous — it compelled them to concentrate on the intellectual labour that machines could not perform. Artificial intelligence is initiating an analogous transformation in the domain of intellectual labour. What it automates is routine intellectual work — information retrieval, material organisation, format adjustment, first-draft generation — and this is precisely what liberates the researcher's time and energy, enabling a focus on the higher-order intellectual work that belongs uniquely to human beings: formulating original research questions, making complex strategic and ethical judgements, and synthesising knowledge across different domains into new insight.

In one project in which the author participated, an explicit institutional arrangement was established to this end. After a rigorous 'steelmanning' process — in which the artificial intelligence scrutinised arguments from multiple opposing perspectives — and a 'triple hostile reader test' — in which arguments were challenged, respectively, from the standpoints of a capitalist critic, a state-socialist critic, and an anarchist critic — the process was compelled to halt before implementing any revisions, awaiting the approval of the human researcher. Without explicit human authorisation, the research could not proceed. The rationale was straightforward: 'To ensure that human beings maintain control over the direction of argumentation.'

The same project also explicitly enumerated those things that artificial intelligence **cannot** replace: theoretical innovation, original analysis, critical thinking, value judgement, and domain expertise. The outputs of artificial intelligence were defined as 'not publication-ready, but an excellent starting point'. This positioning is essential — it acknowledges that artificial intelligence can enormously accelerate the process of knowledge production, while insisting that ultimate intellectual judgement remains with human beings.

By this stage, the AI has produced a well-structured draft — grounded in Amara's documents, organised by her methodology, written in her specified register. Yet the draft lacks something only she can supply. She has spent two years attending parliamentary hearings on mining governance; she has interviewed community leaders in Mtwara who watched gas revenues bypass their region entirely; she knows which civil society organisations have credible data and which reproduce donor talking points. It is Amara who decides that the central argument must foreground the gap between legislated resource nationalism and the contractual reality of existing mining development agreements — a judgement requiring not only information but the political discernment to know what matters most. She rewrites the

recommendations entirely, drawing on conversations no AI has access to, and overrules the AI's cautious 'on the other hand' equivocations where the evidence plainly supports a stronger conclusion.

For researchers in the Global South, this sixth question carries particular political significance. The irreducibility of opinion and insight means that artificial intelligence cannot substitute for knowledge embedded in communities, cannot replace the witness borne to struggle, cannot replicate the understanding that comes from standing alongside the oppressed. For generations, the power over knowledge production has been concentrated in the well-resourced institutions of the Global North. The proliferation of AI tools creates possibilities for resource-constrained institutions in the South — but only on one condition: that researchers maintain a firm hold on decision-making authority. What questions are worth posing, whose side to take, which sources to trust, what methods to employ, to whom to speak — these are choices that artificial intelligence cannot and should not make on your behalf.

The ultimate purpose of mastering knowledge engineering is not to become absorbed in dialogue with machines, but to employ machines to liberate us for the most essential of human tasks — to go deeper into the field, to listen to the 'silent voices', to acquire the tacit knowledge and popular wisdom that cannot be digitised. Genuine knowledge production has always moved 'from the masses, to the masses' — artificial intelligence makes this path more solid to walk.

Conclusion: Six Questions, One Whole

The six questions advanced in this article do not constitute a checklist but an organic whole. Problem orientation shapes the epistemological framework — what you seek to investigate determines from what standpoint you must proceed. The epistemological framework determines which information is relevant — your analytical perspective dictates what data is worth collecting and which sources deserve trust. The information reserve is transformed into analysis through methodology — abundant raw materials require a precise process of elaboration to yield valuable insight. The outputs of methodology are conveyed to audiences through form of expression — the same analytical conclusions require different modes of expression to reach different readers. And opinion and insight — the researcher's distinctive contribution — drives the entire process: it is what determines which questions are worth posing, which standpoints are worth maintaining, and whether the final product carries meaning.

Six weeks later, Amara's report is complete. It is not the product the AI would have generated on its own — that initial output was a generic survey devoid of standpoint or analytical architecture. Nor is it the product she would have produced without the AI — working alone, she could not have mapped the full fiscal architecture, cross-referenced legislative provisions against contractual terms, and produced a polished 8,000-word report in the time available. The report is the product of a structured collaboration: Amara supplied the questions, the standpoint, the documents, the methodology, the formal requirements, and the final judgement; the AI supplied the capacity to process, organise, draft, and redraft at a speed no solo researcher could match. Knowledge engineering made this collaboration productive; without it, the same tool would have produced the same bland overview with which she began.

As researchers systematically apply these six questions, a practical transformation occurs naturally: prompts inevitably grow longer and more complex. When you must simultaneously convey problem orientation, epistemological framework, information reserve, methodological prescriptions, and formal requirements, a single prompt quickly becomes a guidance document running to thousands of words. When a single prompt can no longer accommodate sufficient depth across all six dimensions, the natural next step is to distribute different dimensions of the work across multiple specialised agents — each agent responsible for one dimension, coordinated by orchestration logic. This is precisely the problem that declarative multi-agent systems are designed to address, as we have discussed in detail in a separate article.

Yet the core message of this article concerns not the direction of technological evolution but a more fundamental recognition: knowledge engineering is not a technical skill but a practice of intellectual discipline. It demands that, before asking a machine to speak on your behalf, you first know clearly what you know, what you believe, and what you wish to say. The true value of these six questions lies not in their capacity to help you write better prompts — although they do achieve this — but in the fact that the process of answering them is itself the core of serious knowledge production.

Artificial intelligence has accelerated the process of knowledge production, but it has not altered its essential nature. The value of knowledge continues to derive from deep insight into reality, from the conscious choice of standpoint, from the commitment to standing with those under investigation. To master knowledge engineering is to better accomplish this work that belongs to human beings — not to surrender it to machines.