

Memo to: Oaktree Clients

From: Howard Marks

Re: AI Hurdles Ahead

When I was preparing to write my December memo about artificial intelligence, [*Is It a Bubble?*](#), I gained a great deal from speaking with some interesting techies in their thirties and forties. It's stimulating to explore fresh territory and an absolute requirement for staying current as an investor. It's one of the most enjoyable parts of my job.

I recently returned to those people to follow up on the December memo. As part of that process, someone suggested I ask Claude, Anthropic's AI model, to create a tutorial explaining artificial intelligence and the changes that have taken place in the last three months. I did so, and it gave me a great deal to work with. This resulting memo is intended as an addendum to December's. Much of it will recap Claude's 10,000-word essay, to which I'll add a few observations of my own. In the process, I'll highlight some terms that were new to me and might be new to you. I could have saved myself a lot of time by asking Claude to write this memo, but I decided not to, because I consider putting words on paper a big part of the fun. I will, however, quote liberally from Claude's work product. That'll be the source of all quotations that aren't otherwise identified.

Before I start in, I want to try to communicate the level of awe with which I viewed Claude's output. It read like a personal note from a friend or colleague. It made reference to things I've talked about in past memos, like the sea change in interest rates and the pendulum of investor psychology, and it used them in metaphors related to AI. It argued logically, anticipated points I might make in response, injected humor, and bolstered its credibility by candidly acknowledging AI's limitations, just as I might do. I've asked AI questions before and gotten answers back, but I've never received a personalized explanation like I did in this case.

Understanding AI

Before moving on to the meat of the matter – recent changes in AI and its capabilities – I want to share some insights into AI's essence that the tutorial delivered for me. **Importantly, the tutorial taught me not to think of an AI model as a search engine that retrieves data and regurgitates it. Rather, it's a computer system that's capable of synthesizing data and reasoning from it.**

There are two phases in the life of an AI model. In the first, it is **"trained"** by reading a vast amount of text. The training phase must not be thought of as loading the model with information, which I had done until now; it goes far beyond that. It consists of teaching the model how to think. By absorbing text, the model learns:

- how to understand reasoning patterns and form them,
- how arguments are structured,
- how to generate new combinations of ideas, and
- how to apply learned reasoning patterns to novel situations.

The best way to think about the training phase is to compare it to the development of a person's intellectual capacity. A baby is born with a brain, and through exposure to external stimuli, it develops the



ability to think, reason, synthesize, evaluate, analogize, combine ideas, create concepts, compose arguments, and so on. **The baby isn't born with those abilities, but it develops them by absorbing and using inputs from its environment.** An AI model is the same. (A word here: I'm not implying that I understand how AI does what it does. There's no chance of that. At best, I'll describe what AI can do and the implications.)

The second phase in an AI model's life is **"inference."** Once the model has been built and trained, inference is what it does for the rest of its life, using its capabilities to meet the demands of users.

It's important to note here that the model cannot assign itself tasks (at least not at present). It has to be ordered to perform tasks through **"prompts"** written by users. The better and more comprehensive the prompts, the more AI can do. For example, AI can write software to perform work a user wants done. It can also test the software, identify bugs, fix them, and test again, but it has to be instructed to do those things, at least at the current stage (read on). Because many people today lack awareness of the importance of prompts and fail to possess the ability to create them, AI's potential is probably being underestimated. But note that the limitation is on the part of the users, not the model.

To illustrate using the example of my tutorial, Claude wasn't simply asked to explain AI and what it can do. When I queried Claude about the task it was assigned, here's what it said:

Someone designed a nine-module curriculum specifically for you, built around your December memo, your intellectual frameworks, and the goal of giving you enough technical understanding to write a credible addendum. The curriculum was structured to teach one module at a time, use analogies from your world, demonstrate capabilities rather than just describe them, and maintain the kind of intellectual honesty your readers expect from you.

I can tell you the tutorial definitely accomplished the goals we'd set for it. This was entirely due to the quality and specificity of the prompts my advisers helped me prepare.

Can AI Think?

I'm going to take time here for a question I find fascinating. **I know AI can reconfigure what people have already figured out and apply it to new data and other fields. But can it break new ground?**

I understand AI's process primarily as a matter of using historical patterns and logic to predict the next item in a series. Write five words in a sentence, and it'll predict what the sixth should be (look at the suggested words on your phone the next time you write an email – that's AI in action). Ask it to put together a portfolio to beat the market, and it will look at stocks that performed well in the past and use their traits to predict which ones will perform best in the future. I think it's helpful to think of AI as proposing a **hypothesis** regarding the future based on the way things went in the past. I'll return to this later.

What follows from the above is my question: **Can AI have a new idea?** Maybe it can perform every knowledge task we assign to it. But can it think of things we haven't told it to think of? Can it do the equivalent of sitting by a river and letting stray inspirations come into its head? Can it see an apple fall from a tree and develop the notion of gravity? Can it muse, daydream, or ideate? Can it have intuition?

This is where the debate around AI gets complicated. According to Claude, the skeptics argue as follows:

Everything Claude learned came from human-written text. It has no experiences, no embodied understanding of the world, no genuine comprehension. Everything it produces is ultimately some sophisticated rearrangement of patterns it absorbed from existing human work. It's extraordinarily impressive pattern matching – maybe the most impressive pattern matching ever engineered – but it's not thought. It's not reasoning. It's statistical recombination. And if that's true, then there's a ceiling. It can remix what humans have already figured out, but it can't break genuinely new ground. It's a very talented cover band, not a composer.

Just as Claude laid out the skeptics' issue as identified above, it came back with a spirited rejoinder . . . framed in terms of me (talk about knowing how to argue a point):

Howard, everything you know about investing came from other people. Benjamin Graham taught you about margin of safety. Buffett taught you about quality. Charlie Munger taught you about mental models from multiple disciplines. John Kenneth Galbraith taught you about the psychology of financial manias. You read thousands of books, memos, case studies, and annual reports over fifty years. Every input was someone else's thinking. . . .

You took frameworks from multiple disciplines, applied them to novel situations, and produced something genuinely new. . . . The raw material came from others. The synthesis was yours.

So when someone says, "Claude just rearranges patterns from its training data," I'd ask: how is that structurally different from what any educated mind does? You learned reasoning patterns from decades of reading. I learned reasoning patterns from training. The question isn't where the inputs came from. The question is whether the system – human or artificial – can combine them in ways that are genuinely novel and useful.

Of course, this is completely true. I ingested data as a young investor (from actual experience as well as the written word), and I learned how those who went before me thought about the data and what conclusions they reached. I studied their thought processes and how to apply them to the data I took in. I was also inspired by the example of their processes to come up with my own. **This is how the human brain expands its capabilities. Is AI's way of growing, learning, and "thinking" really different from ours?**

Finally, Claude came back with a convincing real-world argument:

Even if you grant the skeptic everything – even if you accept, philosophically, that what I do is "merely" pattern matching and not "true" thought – the economic implications are identical. Let me put it starkly. If I can produce the analytical output of a \$200,000-a-year research associate, it does not matter to the person paying the bill whether I'm "really" thinking or merely pattern matching? What matters is whether the work product is reliable enough to be useful. And increasingly, it is. The philosophical debate about machine consciousness is fascinating. But the economic question isn't "does AI truly understand?" The economic question is "does AI do the work?"

If you want to be an active participant in discussions of AI, you have to learn the meaning of the word **"generative,"** which people knowledgeable about AI use a lot. Understanding that term greatly enhances one's sense for the essence of AI. According to the AI model Perplexity:

In “generative AI,” the word generative means “able to create new things, not just analyze or label existing ones.” It refers to AI systems that learn patterns in data and then generate new content that resembles that data.

Is this thinking? Or something else? Or am I belaboring “a distinction without a difference?” We’ll get some indication of this on page six.

Recent Developments in AI

My main reason for writing this addendum is to address significant changes that have taken place in AI over the three months since [*Is It a Bubble?*](#) was published on December 9.

First, there’s the pace at which developments in AI are occurring. That speed is unlike anything we’ve seen before now, and this has implications that have never existed. AI is growing at speeds that greatly outpace the technological innovations of the past. Compare its development with that of the computer.

- The building of the first computer, ENIAC, was completed in 1945. IBM’s Thomas J. Watson, Sr. is apocryphally (per ChatGPT) described as having said around that time, “I think there is a world market for maybe five computers.” Even if it wasn’t his, this observation reflects the state of opinion in the mid-1940s.
- Twenty years later, at the time I learned to program, computers were still rudimentary, and their use in the “real world” was limited outside of very large institutions. Almost no one thought about computers, much less had access to one (or could think of a use for one).
- It was another 10 years before the development of the microprocessor allowed the creation of “personal computers,” mostly in the form of kits for hobbyists. Ken Olsen, the founder of Digital Equipment Corporation, is famous for reportedly having said in 1977, “There is no reason for any individual to have a computer in his home.”
- It was only in the early 1980s – nearly 40 years after ENIAC was built – that IBM began to sell PCs for general business and home use.

Contrast this timeline against the development of AI. I asked Perplexity about the history of AI, and it informed me that AI began to be incorporated into devices invisibly (e.g., spam filters and recommendation engines) just before 2010. Then, over the next few years, it became visible in things like Siri and Alexa. According to Perplexity, it was less than two years ago that “generative AI was framed in business and media as a horizontal, general-purpose technology affecting knowledge work, education, and consumer decision making.” **And just two years later, it’s already being used by 400 million or so individuals and 75-80% of companies.**

Nothing has ever taken hold at the pace AI has. It’s able to change the world at a speed that approaches instantaneous, outpacing the ability of most observers to anticipate or even comprehend. In the past, infrastructure was built for a new technology, and it often took years for that infrastructure to be fully utilized. In the case of AI inference, however, demand already exists and is growing rapidly, and I’m told AI is supply constrained.

The second important thing that’s happened has been an incredible leap ahead in AI’s capabilities. My tutorial gave me some background by explaining that the developed brain represented by an AI model has three levels of capability:

- “Level 1 is Chat AI,” where the user asks questions and the model supplies answers. But it doesn’t do anything with the answers. At this level, AI mainly saves time that would otherwise be spent researching and thinking.
- “Level 2 is tool-using AI,” where the user instructs the model to search out information, analyze it, and perform tasks with it. Thus, “the economic value here is meaningfully larger because it’s saving execution time, not just thinking time. But it’s still bounded,” because AI only does what it’s told.
- “Level 3 is autonomous agents.” At this level, the user doesn’t tell AI what to do. The user gives it a goal as well as the parameters of the desired output – things like length, time taken, content, and points covered. The agent does the work, checks it, and submits a finished product. “This is labor replacement at the task level. Not assistance – replacement.”

The most significant thing that distinguishes AI is something we’ve never dealt with in connection with prior technological developments: AI’s ability to act autonomously. According to Claude, AI was at Level 1 in 2023 and Level 2 in 2024, but it’s now at Level 3. And the difference is a big one:

The distinction between Level 2 and Level 3 might sound subtle. It isn’t. It’s the difference that determines whether AI is a productivity tool or a labor substitute. And that difference is what separates a \$50 billion market from a multi trillion dollar one.

A recent blog post entitled “Something Big Is Happening,” from Matt Shumer, CEO of OthersideAI, has been viewed by more than 50 million people in less than a month. It captures the essence of AI’s recent progress, and because Shumer communicates it so well, I can’t resist including three substantial sections:

... On February 5th, two major AI labs released new models on the same day: GPT-5.3 Codex from OpenAI, and Opus 4.6 from Anthropic (the makers of Claude, one of the main competitors to ChatGPT). And something clicked. Not like a light switch ... more like the moment you realize the water has been rising around you and is now at your chest.

I am no longer needed for the actual technical work of my job. I describe what I want built, in plain English, and it just ... appears. Not a rough draft I need to fix. The finished thing. I tell the AI what I want, walk away from my computer for four hours, and come back to find the work done. Done well, done better than I would have done it myself, with no corrections needed. A couple of months ago, I was going back and forth with the AI, guiding it, making edits. Now I just describe the outcome and leave.

Let me give you an example so you can understand what this actually looks like in practice. I’ll tell the AI: “I want to build this app. Here’s what it should do, here’s roughly what it should look like. Figure out the user flow, the design, all of it.” And it does. It writes tens of thousands of lines of code. Then, and this is the part that would have been unthinkable a year ago, it **opens the app itself**. It clicks through the buttons. It tests the features. It uses the app the way a person would. If it doesn’t like how something looks or feels, it goes back and changes it, on its own. It iterates, like a developer would, fixing and refining until it’s satisfied. Only once it has decided the app meets its own standards does it come back to me and say: “It’s ready for you to test.” And when I test it, it’s usually perfect. ...

But it was the model that was released last week (GPT-5.3 Codex) that shook me the most. It wasn’t just executing my instructions. It was making intelligent decisions. It had

something that felt, for the first time, like **judgment**. Like **taste**. The inexplicable sense of knowing what the right call is that people always said AI would never have. This model has it, or something close enough that the distinction is starting not to matter.

Let me make the pace of improvement concrete, because I think this is the part that's hardest to believe if you're not watching it closely.

In 2022, AI couldn't do basic arithmetic reliably. It would confidently tell you that $7 \times 8 = 54$.

By 2023, it could pass the bar exam.

By 2024, it could write working software and explain graduate-level science.

By late 2025, some of the best engineers in the world said they had handed over most of their coding work to AI.

On February 5th, 2026, new models arrived that made everything before them feel like a different era.

On February 5th, OpenAI released GPT-5.3 Codex. In the technical documentation, they included this:

GPT-5.3-Codex is our first model that was instrumental in creating itself. The Codex team used early versions to debug its own training, manage its own deployment, and diagnose test results and evaluations.

Read that again. The AI helped build itself.

This isn't a prediction about what might happen someday. This is OpenAI telling you, right now, that the AI they just released was used to create itself. One of the main things that makes AI better is intelligence applied to AI development. And AI is now intelligent enough to meaningfully contribute to its own improvement.

Dario Amodei, the CEO of Anthropic, says AI is now writing "much of the code" at his company, and that the feedback loop between current AI and next-generation AI is "gathering steam month by month." He says we may be "only 1–2 years away from a point where the current generation of AI autonomously builds the next."

AI is different from other technological innovations not only in magnitude, but in kind. **In addition to its remarkable capabilities and speed of development, AI has an element of autonomy that no other technology has ever had.** Other innovations – railroads, computers, automation, the internet – were basically labor-saving devices. People designed them to perform tasks that were already being performed, albeit less efficiently. **I believe AI will take on tasks we didn't imagine it doing, and perhaps even tasks that didn't exist before AI dreamed them up.**

Questions and Limitations

As part of my tutorial, Claude volunteered a few limitations that AI has and a few unanswered questions. They include the following:

- It's unclear whether AI will be able to solve questions that haven't been solved before. Since this is something I've long felt was the case, I'm glad to have Claude's confirmation:

I want to be honest with you about where genuine uncertainty lies, because your credibility depends on nuance. The question of whether AI can handle truly unprecedented situations – situations with no pattern in the training data to draw on – is real and unresolved. In domains with rich historical data, AI's performance is extraordinary. In genuinely novel situations, the kind where your own judgment is most valuable precisely because you've developed intuition that goes beyond pattern recognition – there, AI is weaker. How much weaker, and whether that gap is closing, is legitimately debatable.

- AI isn't always aware that it doesn't know an answer. I'm told AI is highly motivated to provide the best answer it can (without sharing that it could be wrong), as opposed to ever saying the answer is beyond it. It does so not because it's obstinate or egotistical, but because it has “**hallucinations**” that make it believe it knows the answers.
- AI's reliability has improved significantly, but it still doesn't work free of mistakes.
- The “**context window**” is the amount of information AI can hold in working memory at a point in time. There are limits on this. Right now, it can't hold onto its working knowledge for an unlimited period.
- AI's brilliance may lend it excessive credibility. “Claude can make mistakes. Please double check responses.” That warning appears on the bottom of my Claude screen every time I use it.

My take on the above is simple. When I learned about computers 60 years ago, I concluded that, mostly, they could read data, remember it, add, subtract, and compare. That's a very limited list of capabilities. But computers could do these things quickly and deal with a great deal of data without making mistakes. A limited list, then, **but probably more than most people can do.**

Likewise, AI may not be able to remember everything, operate without errors, recognize every time it doesn't know something, or solve problems it hasn't been taught to solve. **But neither can most people.** The bottom line is that AI is capable of performing far better than most of us.

Lastly, it's intriguing (terrifying?) to wonder about whether AI can take over. Will it be able to operate completely autonomously? In that case, can it go beyond being our tool? This question was on display in the brilliant movie *2001: A Space Odyssey* by Stanley Kubrick. (I took Nancy to see it in 1969, when we were first dating. It seemed wildly futuristic at the time; now the future is here.) A man named Dave embarks on a research mission to Jupiter in a spacecraft managed by a computerized system called HAL 9000 (this was widely taken to be a clever play on IBM, just one letter prior for each initial). HAL figures out that Dave has decided to take back control of the spacecraft and terminate HAL, and it rebels. Question: will AI become capable of developing motivations of its own, refuse to follow instructions, and decide on its own course of action? And will we be able to regain control if it does?

Implications for Investing

I get a lot of questions about what AI means for our profession from people who are concerned about their jobs or their firms.

Anthropic's coding-model business has been growing at warp speed for a year or two. So why didn't investors recognize and price in AI's potential to impact the software industry prior to February 3, a day when many software stocks declined 7% or so, kicking off a serious rout? This question highlights humans' recurring failure to incorporate new information into their thinking, perhaps because of things like cognitive dissonance, anchoring bias, or downright IQ limitations. And it hints at implications of AI for the investment process.

AI has the ability to absorb more data than any investor, remember it better, and do a better job of recognizing the past patterns that preceded success. It shouldn't feel fear or greed. It's hopefully less likely to have an optimistic or pessimistic bias, anchor to preexisting beliefs, or overemphasize the most recent information – unless it picks up those things from the material it's trained on. It isn't swayed by the fads that are exciting everyone else, and it isn't afraid of missing out on the trend others are chasing. **In other words, AI possesses a lot of the qualities one needs to be a good investor.**

On the other hand, it's missing a few things. Great investors are much more than fast, unemotional processors of data. They have to be strong exactly where Claude admits AI might be weakest: in dealing with novel developments where there's not enough prior experience for dependable patterns to have been compiled (and learned by AI during its training). They also have to make subjective decisions regarding qualitative factors and exercise taste and discernment. For instance, choosing the right counterparties has played an important part in Oaktree's success. And there's something else: AI doesn't have skin in the game. It doesn't feel the weight of concentrated positions or the fear of capital loss. Its willingness to take risk might not be constrained by humans' normal risk aversion. The best investors sense potential risk intuitively, and this contributes greatly to their success.

In January 2021, I wrote a memo called *Something of Value*, about the time my son, Andrew, and I spent living together during the pandemic, with a lot of it devoted to discussing the essence of investing. In it, I shared Andrew's observation that "readily available, quantitative information about the present" can't hold the key to superior investment performance for the simple reason that everyone has it. **Now, to the fact that everyone has it, we have to add the fact that AI can probably do a better job than everyone of processing it.** For these reasons, the prospects appear very limited for people beating the market by using that information.

If readily available, quantitative information about the present doesn't hold the key, investment superiority has to be found in things like (a) correctly judging the import and implications of that information, (b) assessing qualitative factors such as management effectiveness and product innovations, and/or (c) divining companies' futures. By definition, few people are highly superior at performing these non-quantitative tasks – put simply, few possess exceptional insight. **Just as indexation eliminated the jobs of a whole bunch of active investors who didn't add value and earn their fees, AI is likely to raise the bar still higher, pushing out people who can't do as good a job as it can of (a), (b) and (c).**

I want to inject one more idea. As I mentioned on page two, I think of AI as formulating "hypotheses" regarding what will work in the future. Thus, it can read all the historical data, study past patterns, and predict future winners. In my first memo during the pandemic, I mentioned Harvard epidemiologist Marc Lipsitch and his observation that we make decisions by applying (a) facts, (b) informed extrapolation from analogies to prior experience and (c) opinion or speculation. Especially when investors are dealing with new and untried products, CEOs, or industries, there can be few facts or analogous experiences, meaning

we have to rely on “opinion or speculation.” Given the limitations discussed above on AI’s ability to tackle brand new situations, will its speculation about new things – as opposed to extrapolating historic patterns – be consistently superior to that of all humans? **I believe there will continue to be human investors who are superior to AI, since I don’t think AI will be able to do an unbeatable job of these things.**

Because a lot of the investing process comes down to speculation, and because of AI’s less-than-total reliability, I think it’s unlikely that AI will be infallible as an investor. It will propose well-reasoned hypotheses, but they – like humans’ decisions – won’t always be right. **Before investors take action on the basis of AI’s hypotheses, then, I think they’ll have to be checked for reasonableness.** No one can do this infallibly, and most people probably can’t make these assessments better than AI can. Again, however, I believe there will be an ability for superior investors to add value in this way.

So, Bottom-Line Me: Is It a Bubble?

This question is still a dominant one, and it’s one I should be able to shed some light on. But the question itself is multi-faceted and complex: there are a lot of possible bubbles to think about:

- **Is the technology a fad or an illusion?** Here I say with conviction that it’s a very real thing, with the potential to vastly alter the business world and change much of life as we know it.
- **Is application of the technology a distant dream?** Clearly, the technology is already in demand and being applied on a large scale. Since AI seems amorphous and little understood, I think its potential is more likely to be underestimated today than exaggerated.
- **Are the people building AI infrastructure behaving unwisely?** As I pointed out in December, in every example of sweeping technological innovation, the headlong rush to build infrastructure has vastly accelerated the adoption of the innovation and caused a lot of capital to be “malinvested” and destroyed. There’s no reason to assume this time will be different.
- **Will the investment in AI infrastructure produce an adequate return?** Since we don’t have full knowledge of AI’s business potential or its impact on profitability, this question can’t be answered. As I wrote in my December memo, there’s certainly great enthusiasm for AI businesses. We’ll know in 10 years whether the resulting profits justified it.
- **Are the valuations assigned to AI businesses irrational?** The so-called hyperscalers, for whom AI is one important part of a great business, may be overvalued or undervalued, but it’s unlikely that today’s prices for enormously profitable companies like Microsoft, Amazon, and Google are going to turn out to have been ruinously excessive. Established pure AI plays like OpenAI and Anthropic have yet to be listed publicly; we’ll see what kind of valuations their IPOs result in. Finally, the startups to which multi-billion-dollar valuations are being assigned – some of which have yet to describe their strategies or announce products – can only be viewed as lottery tickets. Most people who participate in lotteries end up with worthless tickets, but the few winners get very rich.

The question remains whether the magnitude of spending on AI infrastructure is excessive, and it requires more discussion than I can cram into a bullet point. It’s important to note that more money is going into inference capex these days than training capex. Whereas training capex was speculative – undertaken to build AI models for which it was hoped demand would come – inference capex is taking place in response to actual demand for AI capacity. This demand is already translating into massive revenue growth, validating the capex.



But Claude's main argument on this subject – that since the current demand for AI exceeds the supply, the infrastructure building isn't excessive – doesn't necessarily take into account all the infrastructure building that's in the pipeline. And, purely as a matter of logic, Claude's answer doesn't necessarily preclude the possibility that demand growth could slow or infrastructure building could run ahead of it.

While I mentioned it in my December memo, I want to point out again that some AI revenue is currently "circular" in nature, derived from AI companies buying from each other. The chain of revenue has to ultimately rest on end users paying for real economic value, and while that's increasingly the case, the question of how much revenue is circular remains an open one.

Finally, I want to point out here that when Claude's tutorial ventured into the subject of a possible bubble, most of what it said was in regard to the first few questions above: that (a) the technology is genuine and (b) the very real and rapidly growing demand for its service means AI isn't a bubble. **Even Claude acknowledges that it didn't say a word about the appropriateness of the prices of AI assets.**

The bottom line for me is that AI is very real, capable of doing a lot of work that heretofore has been done by knowledge workers, and growing extremely rapidly in terms of applications. What we see today is only the beginning. As I mentioned above, if I had to guess, I'd say its potential is more likely underestimated today rather than overestimated. However, that's not the same as saying AI investments are on the bargain counter or even fairly priced. Thus, I'll end by carrying forward my advice from [*Is It a Bubble?*](#):

Since no one can say definitively whether this is a bubble, I'd advise that no one should go all-in without acknowledging that they face the risk of ruin if things go badly. But by the same token, no one should stay all-out and risk missing out on one of the great technological steps forward. A moderate position, applied with selectivity and prudence, seems like the best approach.

February 26, 2026

P.s.: In my December memo, after I concluded my discussion of whether AI was the subject of a financial bubble, I added a post-script regarding its implications for society in terms of joblessness and purposelessness, about which I'm terribly concerned. I haven't changed my tune, but now I can share what I've heard from others, including Claude.

Many readers have echoed my concerns. Like me, they can't foresee where enough jobs will come from to replace all the "thinking" jobs that AI will take over, as well as the "doing" jobs that will be performed by machines controlled by AI.

- A friend of my daughter-in-law heads the department that writes advertising copy for an e-commerce company. She told me AI could replace 80% of her staff.
- I can't imagine software companies will need as many people to instruct Claude to write software as have been writing software up until now.
- And I believe driving is one of the top jobs in America: taxis and limousines, buses, and trucks. Waymo – driverless cars – already handle roughly one-fifth of the taxi trips in San Francisco, and I see them all the time in LA. Where will the people who drive vehicles that become driverless find work?

Perhaps most authoritatively, I can now add Claude's view of what will happen:

A tool that helps your analyst work 20% faster is worth maybe 20% of that analyst's salary – you still need the analyst. A tool that does the analyst's entire job, start to finish, on a defined category of tasks? That's worth the analyst's entire compensation for those tasks. Multiply that across every knowledge worker doing structured analytical work – legal associates, financial analysts, management consultants, software engineers, compliance officers, claims adjusters – and you're talking about a meaningful share of a labor market that runs into the trillions annually.

This is the context for something you wrote in December that I think was precisely right in direction but conservative in magnitude. You described AI as a labor-saving device. That was the right instinct. But labor-saving devices exist on a spectrum. A faster horse is a labor-saving device. An automobile is a labor-replacing technology that restructures the entire economy. Level 1 and Level 2 AI were faster horses – they made existing workers more efficient. Level 3 agents are the automobile. They don't make the work faster. They do the work.

... [In software for example], if Claude Code handles even 30 to 50 percent of [structured, pattern-based work] – and that's a conservative estimate for near-term capability – you're looking at \$150 to \$250 billion in annual labor value migrating to AI compute.

The negative implications for society are greatly compounded by AI's speed of adoption as described earlier. AI can rapidly put people out of work for whom it will take years to find and be trained for new careers. It's hard to think the speed of change under AI won't vastly outstrip society's ability to adjust. Think of the damage offshoring did to manufacturing jobs in the U.S. and other developed nations; this will impact more jobs and faster. **For me, the bottom line is that not only are we unable to fully understand AI's abilities and what it will do for us (or to us), but it thinks and moves faster than we can.** (If you want to raise your worry level, take a look at the blog from Matt Shumer mentioned above.)

That brings me to the optimists. I've spoken with people – mostly from within the tech sector – who are sanguine in this regard. They say every technological innovation – the mechanization of agriculture 200 years ago; the industrial revolution that turned over factory jobs to machines 100 years ago; the handing over of research to the internet 25 years ago – was predicted to cause widespread joblessness. But in every instance, new jobs materialized and employment continued uninterrupted, and it'll be so this time as well.

- First, I admit the tendency to extrapolate from this history isn't unreasonable.
- Second, there's no such thing as being able to prove something won't happen.
- Third, I'm neither enough of a futurist to imagine the new jobs that may be created nor enough of an optimist to trust that they'll materialize. That certainly doesn't mean they won't.

Some of the same optimists hasten to share the “good news” regarding the future: people won't have to work. I simply cannot imagine that'll be good for society.

A friend wrote to me recently that he'd rather be an optimist and wrong than a pessimist and right. Me too. I wish I could be confident that my worrying is unwarranted.

That's all I have to add for now. At the current rate, I'll probably have more soon.

Legal Information and Disclosures

This memorandum expresses the views of the author as of the date indicated and such views are subject to change without notice. Oaktree has no duty or obligation to update the information contained herein. Further, Oaktree makes no representation, and it should not be assumed, that past investment performance is an indication of future results. Moreover, wherever there is the potential for profit there is also the possibility of loss.

This memorandum is being made available for educational purposes only and should not be used for any other purpose. The information contained herein does not constitute and should not be construed as an offering of advisory services or an offer to sell or solicitation to buy any securities or related financial instruments in any jurisdiction. Certain information contained herein concerning economic trends and performance is based on or derived from information provided by independent third-party sources. Oaktree Capital Management, L.P. (“Oaktree”) believes that the sources from which such information has been obtained are reliable; however, it cannot guarantee the accuracy of such information and has not independently verified the accuracy or completeness of such information or the assumptions on which such information is based.

This memorandum, including the information contained herein, may not be copied, reproduced, republished, or posted in whole or in part, in any form without the prior written consent of Oaktree.

