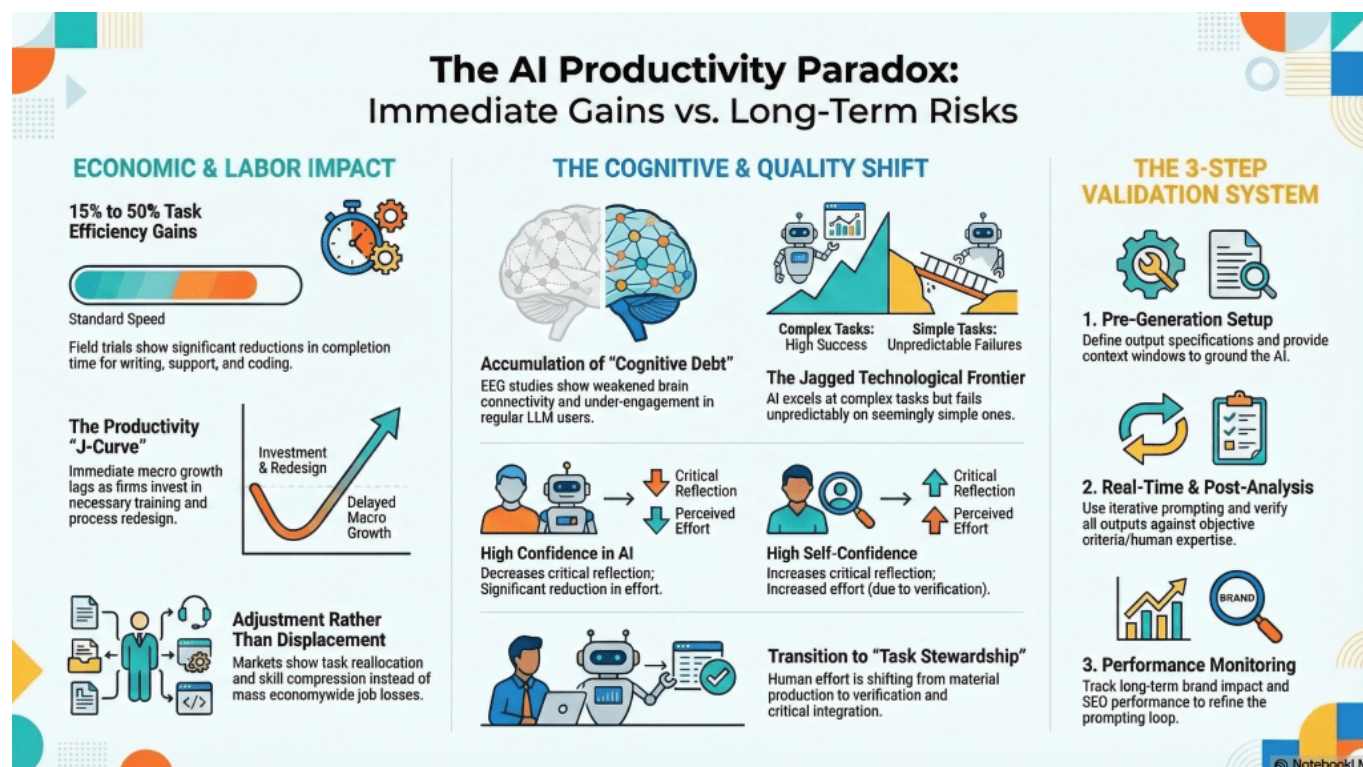


The AI Productivity Paradox: Immediate Gains vs. Long-Term Risks

AI tools are delivering real efficiency wins, but they're also quietly reshaping how workers think, what skills atrophy, and where quality unexpectedly breaks down. Here's what every business leader needs to understand before going all-in.



The AI Productivity Paradox: a framework for understanding short-term efficiency gains alongside emerging cognitive and organizational risks.

There's a quiet tension building inside AI-adopting organizations. On one side: real, measurable productivity gains that no serious executive should dismiss. On the other: a set of slower-moving, harder-to-see risks that, left unmanaged, could erode the very capabilities organizations are counting on AI to amplify.

This tension is what researchers and strategists are calling the **AI Productivity Paradox** and it plays out across three interconnected domains: economic and labor dynamics, cognitive and quality shifts, and the governance frameworks organizations need to navigate both.

The Economic Picture: Real Gains, but Not Instant

Field trials across writing, customer support, and software development consistently show reductions in task completion time of **15% to 50%** compared to standard workflows. That's not marginal, or organizations handling high volumes of routine knowledge work, the compounding effect is substantial.

But those gains don't show up immediately on the macro balance sheet. The **Productivity J-Curve** explains why: in the short term, organizations must absorb the costs of training, workflow redesign, and integration before realizing broader economic returns. Leaders who expect instant ROI are often disappointed, and sometimes abandon AI initiatives right before the curve bends upward.

15-50%

Task Efficiency Gains

Observed across writing, support, and coding workflows in field trials.

J-Curve

Delayed Macro Growth

Short-term investment dip precedes longer-term productivity payoff.

Realloc.

Not Mass Displacement

Labor markets show skill compression and task reallocation, not widespread job loss.

The labor story is similarly nuanced. Rather than triggering the mass displacement many feared, current market data points to **task reallocation and skill compression**, workers shifting away from routine production tasks and toward higher-order judgment, verification, and integration work. The jobs aren't disappearing; they're changing shape.

The Cognitive Risks Nobody Is Talking About Enough

The second domain is where the paradox gets genuinely uncomfortable. Even as AI accelerates output, it may be slowly degrading the underlying human capabilities organizations depend on.

“EEG studies are detecting weakened brain connectivity and reduced cognitive engagement in

regular LLM users, a phenomenon researchers are calling ‘cognitive debt.’”

The mechanism is straightforward: when AI handles the heavy cognitive lifting, such as drafting, reasoning, and synthesis, users engage less deeply with the material. Over time, the neural pathways for critical analysis and creative problem-solving get less exercise. This isn’t theoretical. It’s showing up in neurological data.

There’s also a troubling dynamic around confidence. Research shows that **high confidence in AI output actually reduces critical reflection**; users who trust the tool most are the ones who check it least.

Paradoxically, workers with stronger domain expertise and higher self-confidence engage more critically with AI outputs, applying greater scrutiny and effort to verification. The implication: organizations may want to invest in building genuine expertise rather than assuming AI can substitute for it.

The Jagged Frontier: Where AI Succeeds and Where It Fails

One of the most practically important insights for teams deploying AI is the **Jagged Technological Frontier, as researchers call it**. AI doesn’t fail gradually or predictably; it excels at surprisingly complex tasks, then fails unpredictably on seemingly simple ones.

A system that can draft a sophisticated legal brief may stumble on a straightforward date calculation. A coding assistant that generates elegant architecture may introduce subtle bugs in basic conditional logic. This irregularity makes AI harder to supervise than traditional software, because failure modes don’t follow intuitive patterns. Effective oversight requires humans who understand both the domain and the tool’s specific failure landscape.

Key Terms: A Working Glossary

Glossary of Key Concepts

Cognitive Debt: The gradual erosion of critical thinking and analytical capability that occurs when workers habitually offload complex reasoning to AI. Identified through EEG studies showing reduced brain connectivity in regular LLM users.

The Productivity J-Curve: The pattern where AI adoption initially appears to slow macro productivity growth due to training, integration, and redesign costs before generating compounding returns as workflows mature.

The Jagged Technological Frontier: The uneven capability profile of AI systems, which perform exceptionally well on some complex tasks while failing unpredictably on seemingly simpler ones. Makes AI harder to supervise than traditional tools.

Task Stewardship: The emerging human role in AI-augmented workflows: shifting from direct material production to critical verification, quality integration, and strategic oversight of AI-generated outputs.

Skill Compression: The narrowing of human skill sets observed as AI absorbs routine tasks. Workers increasingly perform a smaller range of higher-level functions, with implications for long-term workforce capability and adaptability.

LLM (Large Language Model): The class of AI systems underlying tools like ChatGPT, Claude, and Gemini. Trained on vast text datasets to generate, analyze, and transform language, the engine powering most current enterprise AI productivity tools.

Pre-Generation Setup: The first step in the 3-Step Validation System: defining output specifications and providing sufficient context before prompting AI, to reduce hallucinations and anchor outputs to accurate information.

Context Window: The amount of text an AI model can “see” and process at once. Providing rich context within this window, such as background documents, specifications, and examples, directly improves output quality and reduces error rates.

A Framework for Sustainable AI Use

The infographic’s 3-Step Validation System offers a practical governance structure that addresses both the quality risks and the cognitive risks simultaneously:

Step 1: Pre-Generation Setup

Define output specifications clearly and load the AI’s context window with grounding information before generating anything. This step dramatically reduces hallucinations and misalignments, and it requires the human to engage meaningfully with the task requirements, counteracting cognitive disengagement.

Step 2: Real-Time & Post-Analysis

Use iterative prompting rather than accepting first outputs, and verify all deliverables against objective criteria or domain expertise. This is where task stewardship happens in practice, and where critical reflection must be deliberately preserved against the pull of over-reliance.

Step 3: Performance Monitoring

Track downstream outcomes, brand impact, SEO performance, error rates, and customer responses to close the feedback loop and continuously refine prompting and verification processes. Organizations that treat AI outputs as the end of the workflow, rather than an input to be refined and measured, will accumulate quality debt they won’t see until it’s costly.

“The organizations that will win with AI aren’t those who use it most; they’re those who’ve built the governance, expertise, and culture to use it best.”

The AI Productivity Paradox isn’t an argument against adopting AI tools. The efficiency gains are real, and the competitive pressure to act is legitimate. It’s an argument for *how* to adopt them: with clear-eyed awareness of the cognitive and quality risks, deliberate governance frameworks, and sustained investment in the human expertise that makes AI outputs actually valuable.

Organizations that manage this balance well will compound both the AI gains and their human capital. Those who don’t will find themselves more efficient at the surface while quietly hollowing out the judgment capabilities they need for anything genuinely difficult.

AI’s \$2 Trillion Moment—and the Hidden Costs We’re Ignoring

Spending on artificial intelligence is expected to cross the \$2 trillion mark by 2026. This massive investment signals that AI is no longer a peripheral experiment but a central part of how global businesses function. Companies are quickly moving past basic chatbots toward agentic systems that can plan and execute complex tasks with very little human help. About 62% of organizations are already testing these autonomous assistants to see how they can improve efficiency. While many people worry about robots taking their jobs, the data suggests a more complicated story. The World Economic Forum predicts that while 92 million roles might disappear by 2030, technology will help create 170 million new ones. This results in a net growth of 78 million jobs, though the transition will likely be quite messy.

For the people actually doing the work, the day-to-day is changing in a major way. We are seeing a shift where knowledge workers move from being creators of content to being stewards of AI systems. This means spending less time on basic execution and more time on verifying and integrating what the AI produces. However, this comes with a strange productivity paradox. Some developers finish their tasks 26% faster with AI, but others actually take 19% longer because they spend so much time fixing mistakes the software made. There is also a real danger of producing what experts call workslop: content that looks good at first glance but lacks any real substance. About 40% of employees have already received this kind of low quality work from colleagues, and it usually takes about two hours to fix each instance.

There are also deeper concerns about what this does to our mental sharpness. A study from the MIT Media Lab suggests that relying too much on AI can lead to cognitive debt, where our brain connectivity actually weakens because we are offloading our thinking. This is particularly true for younger workers, who are seeing a 16% decline in hiring for entry level roles as AI takes over basic tasks. Beyond the human element, businesses are

also struggling with a confusing maze of global rules. The EU AI Act and different state laws in the US often conflict with one another, making it a nightmare for international companies to stay compliant.

Finally, the environmental cost of all this computing power is becoming impossible to ignore. Training just one large model can produce as much carbon as several cars do over their entire lifetimes. This is leading to a new push for Green AI, focusing on energy-efficient hardware such as neuromorphic chips that mimic the human brain. As we head into 2026, the real winners will not be the companies with the most AI, but the ones who can balance speed with high-quality human judgment.

Sources

- [170mn 92mn 78mn – MUFG Americas](#)
- [2026 AI Trends That Will Redefine CIO Strategy – Jade Global](#)
- [AI slop – Wikipedia](#)
- [AI, Productivity, and Labor Markets: A Review of the Empirical Evidence – International Center for Law & Economics](#)
- [An Agentic AI for a New Paradigm in Business Process Development – arXiv](#)
- [Energy-Efficient AI: Green computing approaches for sustainable deep learning – International Journal of Science and Research Archive](#)
- [Energy-Efficient Neuromorphic Computing for Edge AI – arXiv](#)
- [Genesis Mission to Accelerate AI for Scientific Discovery – The White House](#)
- [Global AI Governance & Cross-Border Compliance Risks – Schellman](#)
- [Introducing RAG 2.0 – Contextual AI](#)
- [MAR: Efficient Large Language Models via Module-aware Architecture Refinement – arXiv](#)
- [Quality Control in AI-Produced Content: A Complete Guide – Rellify](#)
- [Small language models \(SLMs\) vs. large language models \(LLMs\) – Invisible Technologies](#)
- [The Great AI Productivity Paradox – DEV Community](#)
- [The Impact of Generative AI on Critical Thinking – Microsoft Research](#)
- [The State of Organizations 2026 – McKinsey](#)
- [The Agentic Organization: A New Operating Model for AI – McKinsey](#)
- [Top AI Trends 2026 – What's Coming Next – Mindrift](#)
- [Your Brain on ChatGPT – MIT Media Lab](#)

From Connectomes to Digital Twins:

Forecasting the Brain in Real Time

Mapping the Living Mind: From Wiring Diagrams to Neural Forecasting

Scientists have spent years trying to figure out how the biological brain works by looking at it from two different angles. One group has focused on connectomics, which is basically mapping the physical wiring of the brain. The other group has looked at functional imaging, or watching neurons fire in real time. We are now seeing these two fields merge through advanced AI to create what researchers call a digital twin of the brain. This move goes beyond just taking high-resolution pictures. It is about building models that can actually predict what a brain will do next.

Building the Physical Maps

The foundation of this work is the wiring diagram. We recently saw a massive milestone with the completion of the central brain connectome for the adult fruit fly, *Drosophila melanogaster*. This map includes more than 125,000 neurons and 50 million synaptic connections. While a fly brain is small, the data is incredibly complex. A single neuron might connect to hundreds of others, making it very difficult to understand how these paths lead to specific behaviors.

We are seeing similar progress in humans too. Researchers recently reconstructed a tiny fragment of the human cerebral cortex. Even though it was only one cubic millimeter in size, it required over a petabyte of data to map at a nanoscale resolution. These physical maps have shown us things we never knew existed, like neurons that form unusual triangular shapes. However, as many experts have pointed out, a connectome is just a map. It does not tell us how the “traffic” of neural activity moves through those wires.

Predicting the Traffic of the Brain

To solve this, researchers are turning to neural forecasting. One of the most important tools in this area is the Zebrafish Activity Prediction Benchmark, or ZAPBench. It uses light sheet microscopy to record the activity of over 70,000 neurons in larval zebrafish. This is currently the only vertebrate where we can see the whole brain active at once at such a high resolution.

By using models originally built for weather forecasting, like those in WeatherBench, scientists are testing how well AI can predict the next 30 seconds of a brain’s activity based on just a few seconds of history. This is a massive shift in how we study neuroscience. Instead of just describing what happened, we are trying to forecast what will happen.

Several new techniques are making this possible:

- **Volumetric Video Models:** Instead of just looking at individual neuron signals, new models like 4D UNets look at the raw 3D video over time. This helps the AI understand the spatial relationships between neurons that other methods might miss.
- **Foundation Models:** Just like the models that power modern chat tools, new foundation models of the mouse visual cortex are being trained on huge amounts of data. These models can be applied to new animals they have never seen before, successfully predicting how their neurons will react to new videos.
- **Classification Strategies:** New architectures like QuantFormer are changing the way we think about brain signals. Instead of trying to predict a continuous wave of activity, they treat neural spikes like a classification problem. This has proven much more effective at capturing the quick, sparse bursts of energy that define how neurons communicate.

Why Global Brain States Matter

One of the biggest hurdles in this research is that a single neuron does not act alone. Its behavior is often influenced by the global state of the brain, such as whether an animal is alert or performing a specific task. A model called POCO, which stands for Population Conditioned forecaster, handles this by looking at local neuron dynamics while also considering the overall state of the entire population. This helps the model understand how shared brain structures influence individual cells.

Future Applications and Interventions

The goal of this research is not just to understand the brain but to interact with it. If we can forecast neural activity in real time, we can develop systems that intervene before something goes wrong. Some models can now run in as little as 3.5 milliseconds. This speed could allow for closed-loop optogenetic interventions, where light is used to stimulate neurons to stop a seizure or a specific craving before the person even realizes it is happening.

We are moving into an era where we can see inside ourselves with the same clarity that we see the world around us. While managing petabytes of data is a major challenge, combining physical maps with AI forecasting brings us much closer to a true mechanistic understanding of intelligence.

This post was written with the help of AI for analysis, using the NotebookLM shared resource here:

<https://notebooklm.google.com/notebook/74dc7f14-54cb-481b-9ee8-8347a6f5cba1>

References and Research Links

- **A Drosophila computational brain model reveals sensorimotor processing**

<https://doi.org/10.1038/s41586-024-07763-9>

- **A connectome and analysis of the adult *Drosophila* central brain** <https://doi.org/10.7554/eLife.57443>
 - **A petavoxel fragment of human cerebral cortex reconstructed at nanoscale resolution**
<https://doi.org/10.1126/science.adk4858>
 - **Foundation model of neural activity predicts response to new stimulus types**
<https://doi.org/10.1038/s41586-025-08829-y>
 - **POCO: Scalable Neural Forecasting through Population Conditioning**
<https://arxiv.org/abs/2410.18025>
 - **QuantFormer: Learning to Quantize for Neural Activity Forecasting**
<https://arxiv.org/abs/2405.17140>
 - **ZAPBench: A Benchmark for Whole-Brain Activity Prediction in Zebrafish**
<https://openreview.net/forum?id=oCHsDpyawq>
 - **Forecasting Whole-Brain Neuronal Activity from Volumetric Video** <https://arxiv.org/abs/2503.00073>
 - **A connectome is not enough - what is still needed to understand the brain of *Drosophila*?**
<https://doi.org/10.1242/jeb.242740>
-

Comparing Artificial Intelligence (AI), Machine Learning (ML) and Deep Learning (DL)

Introduction

Artificial intelligence (AI), machine learning (ML), and deep learning (DL) are terms that are commonly used in the technology industry. While these terms are often used interchangeably, they are not the same thing. Each technology has unique features, advantages, and disadvantages. In this blog post, we will explain the differences between AI, ML, and DL and provide supporting citations from authoritative sources.

Artificial Intelligence (AI) AI refers to the ability of machines to perform tasks that typically require human intelligence. AI is divided into two categories: narrow or weak AI and general or strong AI. Narrow AI is designed to perform specific tasks, such as speech recognition, image recognition, and natural language processing. On the other hand, general AI is designed to perform any intellectual task that a human can do. General AI systems can learn from experience and adapt to new situations. However, as of now, there is no truly general AI in existence, and most AI applications are narrow AI systems.

Machine Learning (ML) ML is a subset of AI that involves training machines to learn from data without being explicitly programmed. In other words, ML algorithms can automatically learn and improve from experience without human intervention. ML algorithms are designed to identify patterns in data and make predictions based on those patterns. The process of training an ML model involves providing it with a large dataset, and then the algorithm will learn to recognize patterns in the data and make predictions based on those patterns.

Deep Learning (DL) DL is a subset of ML that involves training deep neural networks. Neural networks are computing systems inspired by the structure and function of the human brain. These networks consist of layers of interconnected nodes, each of which performs a mathematical operation on the input data. Deep neural networks have multiple layers, which allows them to learn more complex representations of the input data. The process of training a deep neural network involves providing it with a large dataset and adjusting the weights of the nodes to minimize the error between the predicted output and the actual output.

Differences between AI, ML, and DL

Now that we have a basic understanding of AI, ML, and DL, let's take a closer look at the differences between these technologies.

1. Complexity of Tasks

- AI systems are designed to perform tasks that typically require human intelligence, such as speech recognition, image recognition, and natural language processing. ML algorithms are designed to identify patterns in data and make predictions based on those patterns. DL algorithms, on the other hand, are designed to learn from large datasets and can perform tasks that are too complex for traditional ML algorithms. For example, DL algorithms can be used to detect fraud in financial transactions, diagnose medical conditions, and even play complex games such as Go.

2. Type of Learning

- While both ML and DL involve training machines to learn from data, they differ in the type of learning. ML algorithms use supervised, unsupervised, or semi-supervised learning, depending on the problem they are trying to solve. In supervised learning, the algorithm is provided with labeled data, and it learns to make predictions based on that data. In unsupervised learning, the algorithm is provided with unlabeled data, and it learns to identify patterns in the data. In semi-supervised learning, the algorithm is provided with both labeled and unlabeled data. DL algorithms, on the other hand, use a technique called backpropagation to adjust the weights of the nodes in the neural network. This technique involves computing the error between the predicted output and the actual output and then adjusting the weights to minimize that error.

3. Training Data Size

- The size of the training data also differs between AI, ML, and DL. AI systems typically require a large amount of data to learn and perform well. ML algorithms can be trained with smaller datasets than AI systems, but still require a significant amount of data. DL algorithms, however, require large datasets to train deep neural networks. The larger the dataset, the better the performance of the DL algorithm.

4. Hardware Requirements

- DL algorithms require significant computational power and memory to train deep neural networks. As a

result, DL algorithms require specialized hardware such as graphics processing units (GPUs) and tensor processing units (TPUs) to achieve high performance. ML algorithms, on the other hand, can be trained on a standard computer.

5. Interpretability

- Another key difference between AI, ML, and DL is interpretability. AI systems are typically rule-based, and the rules can be easily understood by humans. ML algorithms can be more difficult to interpret, as they learn from data and do not necessarily follow explicit rules. DL algorithms are even more difficult to interpret, as deep neural networks can have millions of parameters and can learn complex relationships between the input and output data.

Conclusion

In conclusion, AI, ML, and DL are all related but distinct technologies with unique features, advantages, and disadvantages. AI refers to machines that can perform tasks that typically require human intelligence, while ML involves training machines to learn from data without being explicitly programmed. DL is a subset of ML that involves training deep neural networks. The key differences between these technologies are the complexity of tasks they can perform, the type of learning they use, the size of the training data, the hardware requirements, and the interpretability of the results. As AI, ML, and DL continue to evolve, they will play an increasingly important role in many aspects of our lives, from healthcare to finance to entertainment.

References:

- Goodfellow, I., Bengio, Y., & Courville, A. (2016). Deep learning. MIT press.
 - Jordan, M. I., & Mitchell, T. M. (2015). Machine learning: Trends, perspectives, and prospects. Science, 349(6245), 255-260.
 - Kelleher, J. D., Tierney, B., & Tierney, B. (2018). Data science: An introduction. CRC Press.
 - LeCun, Y., Bengio, Y., & Hinton, G. (2015). Deep learning. Nature, 521(7553), 436-444.
 - McCarthy, J., Minsky, M. L., Rochester, N., & Shannon, C. E. (2006). A proposal for the Dartmouth summer research project on artificial intelligence, August 31, 1955. AI magazine, 27(4), 12-14.
-

Algorithms for decision making: Free book download from MIT

MIT press has provided a free book on Algorithms for decision making. You can [download it from MIT Press here](#), or alternatively it is available from [this site](#) if the original link fails.

From the data science website:

The book takes an agent based approach

An *agent* is an entity that acts based on observations of its environment. Agents may be physical entities, like humans or robots, or they may be nonphysical entities, such as decision support systems that are implemented entirely in software. The interaction between the agent and the environment follows an *observe-act cycle* or *loop*.

- The agent at time t receives an *observation* of the environment
- Observations are often incomplete or noisy;
- Based in the inputs, the agent then chooses an action at through some decision process.
- This action, such as sounding an alert, may have a nondeterministic effect on the environment.
- The book focusses on agents that interact intelligently to achieve their objectives over time.
- Given the past sequence of observations and knowledge about the environment, the agent must choose an action at that best achieves its objectives in the presence of various sources of uncertainty including:
 1. *outcome uncertainty*, where the effects of our actions are uncertain,
 2. *model uncertainty*, where our model of the problem is uncertain,
 3. *state uncertainty*, where the true state of the environment is uncertain, and
 4. *interaction uncertainty*, where the behavior of the other agents interacting in the environment is uncertain.

The book is organized around these four sources of uncertainty.

Making decisions in the presence of uncertainty is central to the field of *artificial intelligence*

AI use cases - broad applicability

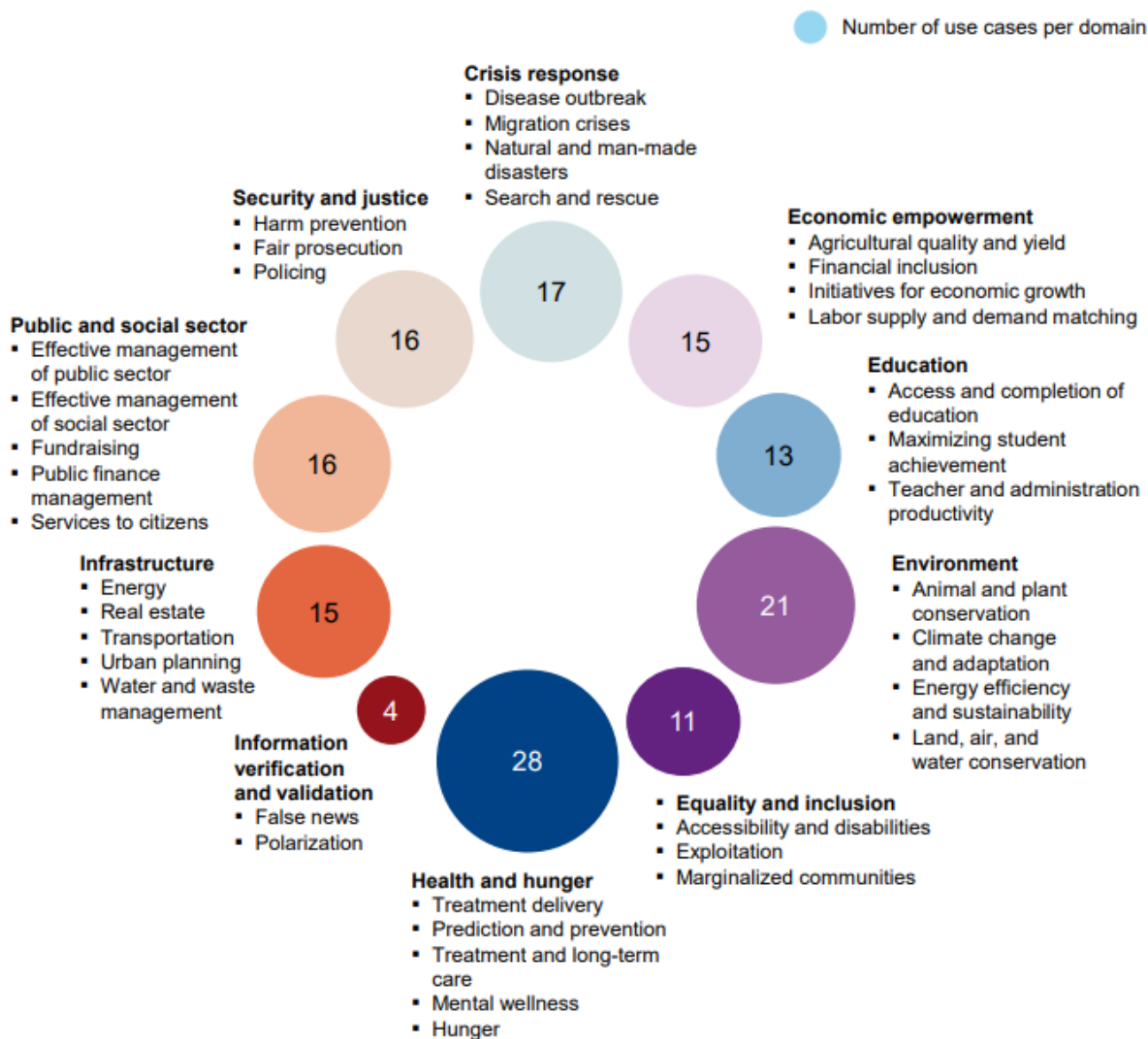
I recently read a paper from the McKinsey global institute, titled "NOTES FROM THE AI FRONTIER, APPLYING AI

FOR SOCIAL GOOD". (Linked below)

The paper defines AI as "Deep Learning" for the purpose of categorization on this particular research. While the definition can be debated, the topic at hand is about the ability to alleviate human suffering, even a bit, using these emerging and growing technologies.



Mapping domains to issue types and use cases in our library.



SOURCE: McKinsey Global Institute analysis

The paper covers potential uses of AI, as well as applied examples, as diverse as AI on satellite maps to predict progression of fires and optimizing responses to managing wildlife poaching monitoring & responses.

Several AI capabilities, primarily in the categories of computer vision and natural language processing, are especially applicable to a wide range of societal challenges. As in the commercial sector, these capabilities are good at **recognizing patterns from the types of data** they use, particularly unstructured data rich in information, such as images, video, and text, and they are

particularly effective at **completing classification and prediction tasks**. Structured deep learning, which applies deep learning techniques to traditional tabular data, is a third AI capability that has broad potential uses for social good. Deep learning applied to structured data can provide advantages over other analytical techniques because it can automate basic feature engineering and can be applied despite lower levels of domain expertise.

MCKINSEY GLOBAL INSTITUTE ANALYSIS (LINKED BELOW)

I think of this applicability in farming and agriculture planning, predicting yields and recommending crop rotation patterns, fertilization based on yield, soil composition, long range forecasting, past local yield data and more. This is something that is completely inaccessible to the general subsistence farmer, but parts of which are regularly employed on larger commercial farming operations in the United States as well as elsewhere. I use this as one example but the potential is amazingly broad.

The challenge lies in bringing these innovative ideas to those who can benefit from them, but who cannot afford to participate in the process as a driver, or sponsor. I am intrigued by the idea of building a collaborative relationship with people living close to the land, directly benefiting from the applicability of these use cases. I picture a supervised learning model, with the on site team being the local populations of the target environments, trained to provide feedback and exploit the findings. The people involved would then collaborate to provide rich on site data, backing up, or reinforcing the model learnings to make it even more capable, backed by the real world data. This is an approach that has boundaries, of course, but would help provide value to a machine learning process, while improving quality of life and as importantly, ensuring that those on the receiving end of this are both the people on the ground and also the technology sponsors – **pride in work, means ownership of progress**. This is a crucial element in helping others, and in my opinion too often lost in simple “aid package” approaches.

Source Paper from Mckinsey:

<https://www.mckinsey.com/~media/mckinsey/featured%20insights/artificial%20intelligence/applying%20artificial%20intelligence%20for%20social%20good/mgi-applying-ai-for-social-good-discussion-paper-dec-2018.ashx>

Also [linked from my blog](#) in case the main link is removed, but [please visit Mckinsey for the original](#) and a great post on the topic.

Getting started with Robotic Process

Automation (RPA)

This post covers the same ground as the link in the papers section of the site, but I think it is worth dropping into a post as well.

Robotic Process Automation brings together classic Business Process Management (BPM) approaches, and macro or scripted automation, with a simpler approach to the automation development. In simple terms, RPA is about making the automation accessible without deep development-based integration.

RPA allows the replacement of repeatable human functions with automation, reducing both costs and time with the assumption that the process is clear, repeatable and well documented or articulated. In addition to simple replacement of human keyboard time, RPA facilitates bringing disparate systems together with minimal direct IT investment, or rigid development. This just touches on the capabilities made accessible by RPA, as it really becomes exponentially more valuable when paired with more advanced technologies such as NLP (Natural Language Processing) on the inbound side, ML (Machine Learning) and advanced analysis / rules engines on the processing side and NLG (Natural Language Generation) on the output side to mimic human to human interaction where that is needed. Think personalization without the cost of the person in the middle!

I drafted the included paper as a short reference on the lessons learned in an early implementation I was a part of, and these lessons continue to hold true as I dive deeper into this topic. I would welcome thoughts on the ideas [outlined in the included paper](#).

Getting Started With RPA

Robotic Process Automation brings together classic Business Process Management (BPM) approaches, and macro or scripted automation, with a simpler approach to the automation development. In simple terms, RPA is about making the automation accessible without deep development-based integration.

RPA allows the replacement of repeatable human functions with automation, reducing both costs and time with the assumption that the process is clear, repeatable and well documented or articulated. In addition to simple replacement of human keyboard time, RPA facilitates bringing disparate systems together with minimal direct IT investment, or rigid development. This just touches on the capabilities made accessible by RPA, as it really becomes exponentially more valuable when paired with more advanced technologies such as NLP (Natural Language Processing) on the inbound side, ML (Machine Learning) and advanced analysis / rules engines on the processing side and NLG (Natural Language Generation) on the output side to mimic human to human

interaction where that is needed. Think personalization without the cost of the person in the middle!

I drafted the included paper as a short reference on the lessons learned in an early implementation I was a part of, and these lessons continue to hold true as I dive deeper into this topic. I would welcome thoughts on the ideas outlined in the included paper. View the paper [here](#).