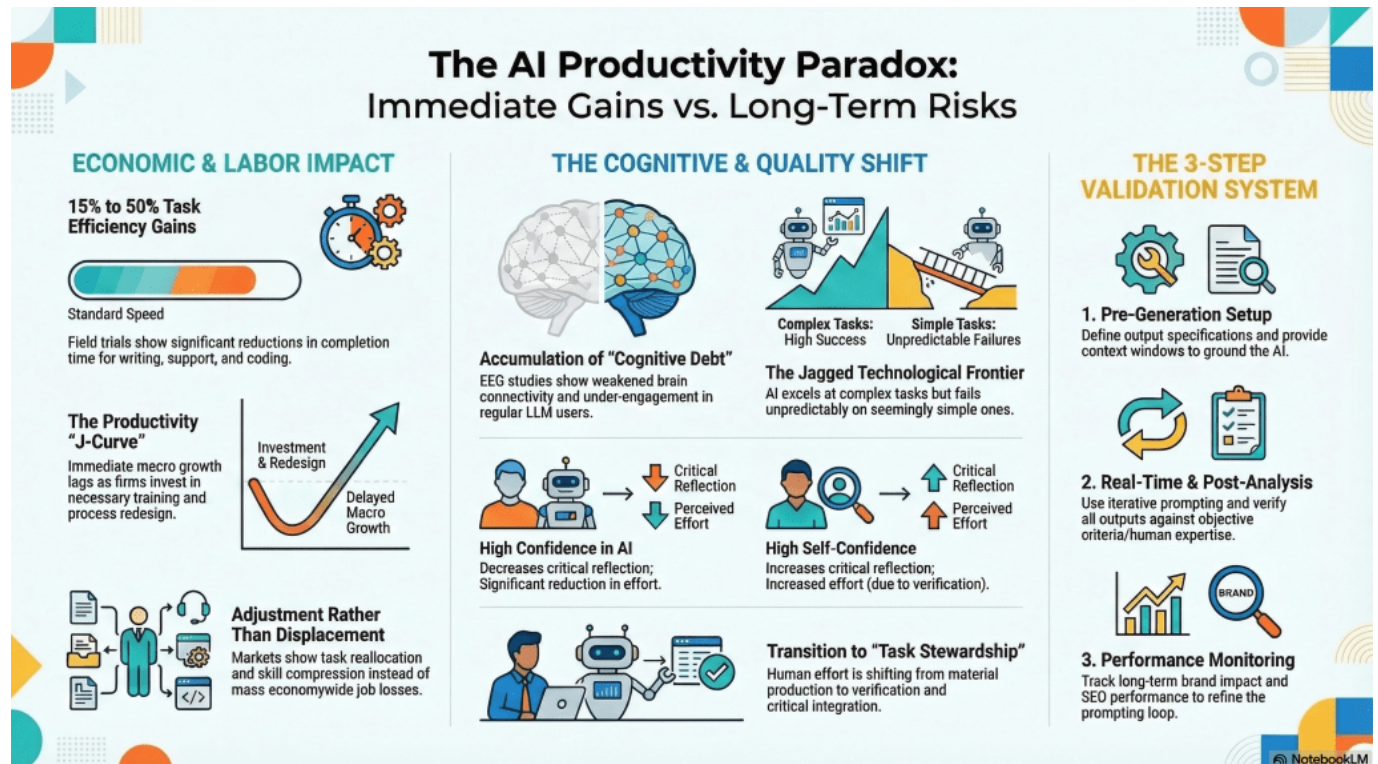


The AI Productivity Paradox: Immediate Gains vs. Long-Term Risks

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.
-

The Best Leaders Have a Contagious Positive Energy

This post references the [HBR article](#) titled “The Best Leaders Have a Contagious Positive Energy” by Emma Seppälä and Kim Cameron.

Take a few minutes to read the whole article [here](#), but one of the key takeaways for me in the value of emotional intelligence and empathy, informing your engaged leadership style. We are all hungry for leaders who care and have a positive energy, you see it in high performing teams where there is an associated high degree of trust. The effort required to project energy and enthusiasm is well worth the investment, but it must be authentic – not the cheerleader style that is empty of real engagement.

Energizers’ greatest secret is that, by uplifting others through authentic, values-based leadership, they end up lifting up both themselves and their organizations. Positive energizers demonstrate and cultivate virtuous actions, including forgiveness, compassion, humility, kindness, trust, integrity, honesty, generosity, gratitude, and recognition in the organization. As a result, everyone flourishes.

[HBR - THE BEST LEADERS HAVE A CONTAGIOUS POSITIVE ENERGY](#)

Digital Twin - exploring the basics

The concept of digital twins is not new, but rather built on ideas that have been explored for the last couple of decades. The technology (compute power, data management & analytics, etc..) and thinking (increasing regulatory and community acceptance of digital approaches to science) have finally hit an inflection point that makes in silico modeling attainable in a cost effective manner.

What this now unlocks is a new opportunity set in the form of machine accessible data, as well as integration of the data sets / ontologies across the target systems / interactions. The need to get to a standardized mechanism to make these data available is tied to the FAIR Data work, and an important dimension to Digital Twin.

Digital twins vs. simulations

Although simulations and digital twins both utilize digital models to replicate a system's various processes, a digital twin is actually a virtual environment, which makes it considerably richer for study. The difference between digital twin and simulation is largely a matter of scale: While a simulation typically studies one particular process, a digital twin can itself run any number of useful simulations in order to study multiple processes.

SOURCE: [IBM](#) , [WHAT IS A DIGITAL TWIN](#)

At its heart, the idea of a digital twin is to reproduce a system in a "runnable" computer model. This oversimplifies the idea, but is a useful construct to think about the problem space and the opportunity it presents. If you can take a scientific instrument, and fully model it in silico, you can then run data sets through it virtually – this makes the assumption that both the inbound and outbound data are available in a machine usable format – something that is tied to this work.

Digital twin is an interdisciplinary research field which includes engineering, computer science, automation and control, and so on. But due to the multidisciplinary nature of the field, it also touches on materials science, communication, operations management, robotics, medicine and other disciplines. A keyword analysis indicates that digital twin, 'smart manufacturing', 'big data', 'cyber-physical system', and 'digital economy' are closely related fields.

SOURCE: "INNOVATIONS IN DIGITAL TWIN RESEARCH" FROM [NATURE PORTFOLIO](#)

The article in nature.com is an interesting piece in that it ties together the many dimensions in this field of research. We can't think of "Digital Twin" as a single entity opportunity, rather to fully realize the potential, we need to look at it as a part of an emerging "virtual capability ecosystem" with applications back to the real world. The value is realized in lower long term costs with increased innovation driven by reduced cost and cycle times, accompanied by increases in application of AI / ML on these models to gain targeted insights that more sharply focus the bench work.

Track the past and help predict the future of any connected environment

SOURCE: [AZURE DIGITAL TWINS](#)

The ability to create learning models for these Digital Twins will improve the accuracy and usefulness of the models over time, and that feedback loop will be a critical part of design. While the industry is maturing, we are seeing more vendors coming to the table with solutions in this space. One of the interesting things to watch is how we as an industry continue to drive open standards in support of these ideas to avoid the traps of "vendor lock in" that were so prevalent in the past.

Books to read: Algorithms to Live By

This book is a solid read with ideas that apply to decision making across a broad spectrum of areas. The authors are able to make the math and conversation around algorithms map to life in well thought and articulated examples that should open your thinking to new ways to approach problems and opportunities.

A few sections that jumped out to me are referenced here or in the reviews, but I encourage you to take the book for a spin yourself.

The most prevalent critique of modern communications is that we are always connected; we're not. The problem is that we are always buffered. The difference is enormous.

ALGORITHMS TO LIVE BY PP226

We are now consuming so much information, we cannot possibly process it all. We now queue information to consume, inhibiting real time engagement and leaving an inescapable feeling of "missing out" or need to "catch up".

From Amazon:

Algorithms to Live By



The COMPUTER SCIENCE of HUMAN DECISIONS

Brian Christian and Tom Griffiths

An exploration of how computer algorithms can be applied to our everyday lives to solve common decision-making problems and illuminate the workings of the human mind.

What should we do, or leave undone, in a day or a lifetime? How much messiness should we accept? What balance of the new and familiar is the most fulfilling? These may seem like uniquely human quandaries, but they are not. Computers, like us, confront limited space and time, so computer scientists have been grappling with similar problems for decades. And the solutions they've found have much to teach us.

In a dazzlingly interdisciplinary work, Brian Christian and Tom Griffiths show how algorithms developed for computers also untangle very human questions. They explain how to have better hunches and when to leave things to chance, how to deal with overwhelming choices and how best to connect with others. From finding a spouse to finding a parking spot, from organizing one's inbox to peering into the future, *Algorithms to Live By* transforms the wisdom of computer science into strategies for human living.

[HTTPS://WWW.AMAZON.COM/ALGORITHMS-LIVE-COMPUTER-SCIENCE-DECISIONS/DP/1627790365](https://www.amazon.com/Algorithms-to-Live-Computer-Science-DECISIONS/DP/1627790365)

Why I recommend this book:

I lead teams in the Pharma / BioPharma industries and we grapple with large challenges on a regular basis – the ideas presented in the book resonate with me as I think about both the scientific / math applications, but as importantly, the human implications. As leaders it is often required that we know enough about everything in our area of responsibility to help guide the decision making for the organization(s). How we go about prioritizing

what to focus on, what to allow in the queue vs what we allow to drop off is a critical bit of surviving and thriving. The best bets are made by those who can separate the noise from the actionable data. To get there, we need to filter and extrapolate from what we have to what we need to do. This book helps shape a number of interesting and workable ideas to explore in this space.

WHY THE PAST 10 YEARS OF AMERICAN LIFE HAVE BEEN UNIQUELY STUPID

I came across this article on Twitter this week and was struck by many of the points. The idea that as we have evolved our social media platforms, we have empowered the worst of society, and amplified their behaviors has become increasingly evident. Read the original article [here](https://www.theatlantic.com/magazine/archive/2022/05/social-media-democracy-trust-babel/629369/):

<https://www.theatlantic.com/magazine/archive/2022/05/social-media-democracy-trust-babel/629369/>

It's been clear for quite a while now that [red America](#) and [blue America](#) are becoming like two different countries claiming the same territory, with two different versions of the Constitution, economics, and American history. But Babel is not a story about tribalism; it's a story about the fragmentation of everything. It's about the shattering of all that had seemed solid, the scattering of people who had been a community. It's a metaphor for what is happening not only between red and blue, but within the left and within the right, as well as within universities, companies, professional associations, museums, and even families.

FROM THE DECEMBER 2001 ISSUE: DAVID BROOKS ON RED AND BLUE AMERICA

Additional Excerpts:

By 2013, social media had become a new game, with dynamics unlike those in 2008. If you were skillful or lucky, you might create a post that would "go viral" and make you "internet famous" for a few days. If you blundered, you could find yourself buried in hateful comments. Your posts rode to fame or ignominy based on the clicks of thousands of strangers, and you in turn contributed thousands of clicks to the game.

This new game encouraged [dishonesty](#) and mob dynamics: Users were guided not just by their true preferences but by their past experiences of reward and punishment, and their prediction of how others would react to each new action. One of the engineers at Twitter who had worked on the "Retweet" button later revealed that he regretted his contribution because it had made Twitter a nastier place. As he watched Twitter mobs forming through the use of the new

tool, [he thought to himself](#), “We might have just handed a 4-year-old a loaded weapon.”

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. *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*

Pistoia Alliance: Patient Centricity

There is an increasing recognition of the value in patient engagement with respect to healthcare in general, as well as the emerging field of personalized / targeted medicine and digital health. The wearable / therapeutic combination, CAR-T therapies, telehealth and so much more fall into this broad category of patient centricity and experience, as well as the direct marketing side of it.

The
Pist
oia
Allia
nce
has
call
ed
for
life
scie
nce
and
heal
thca
re
to
urg
entl
y
rest
ruct
ure
aro
und
pati
ent
cent
ricit
y -
rea
d
the
post
fro



m
the
allia
nce
her
e.

the pandemic has changed behaviors. Billions of people changed the way they interact with healthcare in a matter of months. In this new era of targeted precision medicine, we all play a role in creating the patient-centric future that patients deserve.”

CRISTINA ORTEGA DURAN, CHIEF DIGITAL HEALTH OFFICER R&D
FOR **ASTRAZENECA**

I am excited to see where this leads us as an industry, and how we shift from traditional approaches to include our broad patient populations in developing and delivering medicines and treatments. It will be great to see growing inclusivity across geographic and social boundaries as we increase reach and engagement.