

Dear Investors,

In Q4 2025, Luna returned -8.5% on net assets under management. This brought our 2025 performance to 12.6%. In Q1 2026, Luna returned -24.6% on net assets under management.¹

Artificial Intelligence

Over the last several years, artificial intelligence has become the most prominent thematic issue for investors and the markets. Whether you believe this is just another garden-variety tech cycle in a valuation bubble and/or the end of the modern workforce as we know it, we're undeniably facing what could be a monumental shift in the way we process information and approach the relationship between human and machine intelligence.

As long-term investors, we are guided by the ethos of "strong opinions, weakly held." We build our investment viewpoints with high conviction backed by deep diligence, but with the curiosity and humility to challenge our own assumptions as environments evolve. While the current tech landscape is fast moving and fluid, we see parallels to historical patterns and cycles that could be playing out again.

Below are our current thoughts on the AI tech cycle, with recognition that this, along with frankly most of our work, is one long continuous and never-ending exploration – open to being tested, supported or even disproven. For those less familiar with the space, we've included a brief AI primer.

AI Primer

Key Terms

At its core, AI is the attempt to replicate human intelligence in computer systems. The primary question is how? Classical AI, which was founded in the 1950's, is rule-based and also known as symbolic AI due to its usage of human-readable symbols and logic. This was incredibly difficult and expensive to scale and maintain (in simplistic terms, imagine trying to create if / then statements for everything you wanted a model to be able to do). In contrast, machine learning AI, the relevant branch to today's major tech cycle, is based on pattern recognition and brute force. The basic idea: a machine is fed tons of data; it uses that data to begin to recognize patterns; it receives feedback on those patterns, incorporates that feedback, and refines the pattern further. Ironic to most people's preconceptions on what intelligence is, machine learning is not about logic or reasoning, but

¹ Net asset value calculations are provided by Interactive Brokers and returns are after expenses and fees. Clients should always check their individual statements for returns which may vary due to timing and other factors. Past performance is not indicative of future returns.

rather it is more akin to repeating a question / solution feedback loop trillions of times and incrementally refining it each time until it gets to a generally accepted correct answer (scale is key).

Deep learning is a type of machine learning using neural networks, inspired by the human brain. The transformer architecture is a type of neural network architecture that is foundational to the large language models that are popular today (it enabled parallel processing of text and images). However, most AI used by enterprises today is actually shallow machine learning. Shallow ML uses simpler models where a human must define the relevant input variables. It's most effective at taking in structured data, developing relationships between those pieces of data, and creating predictions based on those relationships (e.g., use cases include credit risk, fraud detection, underwriting, churn / attrition forecasts etc.).

Generative AI (genAI) is deep learning that creates content, but this is relatively new to mainstream usage (within the last few years). Previously, deep learning was primarily focused on analyzing and interpreting existing data (e.g., image classification, face / speech recognition, recommendation systems, fraud detection etc.). Large language models are text-based genAI (e.g., GPT-4) and are the underlying models behind consumer applications like ChatGPT, Claude, Gemini, etc.

Quick side note: for ease, unless specified, when referring to AI or a model generically, we mean deep learning.

Basic Mechanics of a Model

At its core, a model is a large set of numbers. These numbers define relationships between pieces of data. When a question is posed, the model chooses an answer based on what those relationships between the data imply is the most likely correct answer. How did the model get this large set of numbers? As mentioned previously, most of the large models today are based on the transformer architecture. In laymen's terms, programmers create code (aka the transformer architecture) which is like the structure of a brain, including space for a large set of numbers – which you can think of like memory or what a human would draw upon to respond to a question. It's important to note that initially there is nothing in the memory – there's a space for numbers but nothing there yet. The transformer architecture is simply an empty structure or blueprint. Separately, programmers create a training tool. Data is required for the training tool. The training tool is used to ask the brain questions, and the tool gives feedback on the brain's answer. At first, the answers do not make sense since as noted, there is no memory and the numbers are still being created. Each time the brain receives feedback, the brain adds to and refines the large set of numbers. Repeat this process trillions of times, and you eventually get a large set of numbers that reflect fairly accurately relationships between pieces of data. Two important considerations to

note: first, these models are based on probability – based on the large set of numbers defining data relationships, they will give you the answer they believe has the highest probability of being “accurate” (which obviously leaves room for error vs. a deterministic system). Second, “accuracy” is defined by the training tool and its goals, which allows for potential biases.

Key Parts of a Model & Some Quick History

An AI model is composed of four primary parts: processing, model (architecture / training), data, application.

As mentioned, in the training process, models iterate with feedback over and over again. This takes tremendous processing power. This is where the current, incredibly high demand for chips comes into play. There are two main categories of chips: logic and memory. Memory is important, but the most relevant to AI demand right now are the logic chips. Logic chips process information and execute decisions. There are CPUs (central processing units), GPUs (graphic processing units), and NPUs (neural processing units). CPUs historically are the powerhouses – they are the ones executing most tasks in a computer. GPUs were created for displaying graphics (images / videos) and are very important in gaming, for example. NPUs, aka inference chips, were created specifically for AI but primarily for after the training phase; they are good at navigating a neural network after it’s already been developed and trained.

Prior to the early 2010s, AI models were primarily developed and trained using CPUs as the primary processor. However, this was incredibly slow and expensive. In many ways, CPUs were more complicated than AI models required. AI is simply brute force repetition whereas CPUs were set up for more complicated rules and processing. However, in 2012, at the ImageNet Large Scale Visual Recognition Challenge (a competition to evaluate models’ ability to analyze and classify images), the winning model, by a significant margin, was a neural network one called AlexNet. Amongst other innovative developments, AlexNet used GPUs for training. Often considered the “Big Bang” moment for AI, this model opened the door to GPUs becoming the standard chip for machine learning. People realized the graphics use case is very similar to AI training – repeating fairly simple math / iterations over and over again – no need for the complexity of CPUs (and the associated higher time and expense). Another important year to note is 2007 when CUDA (Compute Unified Device Architecture) was released. CUDA is a proprietary program that allows developers to use GPUs for non-graphic purposes (developed by NVIDIA). Obviously, AlexNet and the work leading up to it would not have been possible without CUDA.

Next, you need the right model to wield the processing power effectively. We already mentioned transformers, which is the dominant architecture in the industry today. This was actually released in 2017 by Google in a research paper called “Attention is All You Need” – sometimes referred to as the second “Big Bang” moment in AI. What made this architecture different was how scalable it was. Models based on this architecture could hold and maintain information (without breaking) at levels not previously seen.

Next, to build out the large set of numbers within the architecture (or if using our metaphor from earlier, the memory of the brain), you need the training tool. The training tool is how the model determines what is “right” or “wrong.” It’s important to note here, the company or entity that creates the training tool controls the parameters of this. In the wrong hands, you could see censorship, biases (frankly, this is always likely to an extent, regardless of intention), and so forth.

Hand in hand with training is data. The more data, the more refined a model can be. The models we have today are only possible due to the exponential rise of easily accessible data from increased data storage and internet availability over the last few decades.

The final component is the application – how the consumer or end user ultimately uses the model. ChatGPT, Claude and Gemini have all entered the mainstream and have consumer products that are easily accessible by phone / computer with internet access.

Through looking at the components of a model, we can also understand in simple terms why now? What led to the rise of AI in recent years (and why will it probably continue)? The switch to GPUs allowed for cheaper and faster processing power in training models. A neural network architecture that scaled well was developed and released publicly. The last few decades have seen exponential data growth, aka training fodder.

The next question is how does this landscape play out for the next several years and decades, and then as investors, how do we layer those views onto our investment framework? For this, we look to history.

Layers of the Tech Stack

There are many ways to cut the tech stack. Here is how we think about it – specifically for the purposes of understanding where we could invest along it and how we think value will accrue over time. We’ve also chosen these categories due to their applicability across historical analogues.

At the bottom, there is the infrastructure layer. For AI, this is where we would place everything required for the computing or processing power. This means direct hardware like chips (e.g., NVDA, TSM) and the cloud infrastructure (e.g., AWS, Azure).

Next is the platform layer. This is where the processing power is being harnessed into intelligence. These are the companies that are creating and operating the models (e.g., OpenAI, Anthropic), as well as the companies that have built up around them that make the models more accessible to developers (e.g., PANW – security, SNOW – data / enterprise infrastructure).

At the top is the application layer. This is what a consumer or end user would most likely interface with. It could be an application directly created by a platform company (e.g., ChatGPT by OpenAI) or a separate software company that has incorporated AI into its product offering. Within software, there are both horizontal (e.g., NOW, WDAY) and vertical players (e.g., VEEV, GWRE).

Data obviously plays an important role across the whole tech stack, but generally speaking, it increases in criticality and proprietary value as you move up the stack.

Our core industry thesis is that true long-term compounders will emerge within the application layer, where companies own the mission-critical workflows, proprietary data, and customer relationship. The majority of our investment focus in this space has been on identifying those companies.

In the following paragraphs, we'll walk through how we came to this conclusion – first through historical analogues, followed by comparisons to today's tech cycle, and finally the reality of what we've seen over the last few years.

Some Tech Cycle History

The Mainframe & PC Eras

The 1960's started the mainframe era. Computers were extremely large and expensive, and the industry was dominated by one company, IBM. For the most part, only large enterprises and the government had computers. At the beginning of this era, hardware was critical, and software was simply bundled together with it. Functionally, IBM was the infrastructure, platform and application layer. However, in the late 1960's, IBM, for antitrust reasons, was forced to separate its software and hardware businesses. Additionally, following the prediction of Moore's law (first predicted by Gordon Moore in 1965), compute costs began to dramatically decline. Hardware began to commoditize as we entered into the PC era (1980's on).

At this point, there were several manufacturers including Intel, AMD and Dell. Intel dominated early, but the PC hardware industry has since largely commoditized. The platform layer (for this era, operating systems) saw the rise of Microsoft Windows as well as a few smaller players. Windows continues to be the dominant operating system today (followed by Mac OS), but generally speaking, it quickly became a mature market with its value having faded in relevance with the rise of the internet. Separately packaged software (application layer) such as Microsoft Office and Adobe also popped up during this time. These key applications saw strong, long-term growth and continue to remain highly relevant even today.

Rise of the Internet

The internet started to rise in prominence in the mid-1990's. The Telecommunications Act of 1996 (deregulated the industry allowing for new entrants), combined with strong internet traffic growth, led to huge amounts of capital pouring into all layers of the tech stack. In infrastructure, the investment was focused on fiber-optic networks. There were the network owners, e.g., AT&T, who unfortunately, post Telecom Act, faced significant competition, massive capex (funded primarily with debt), and limited pricing power. The hardware and network software vendors, e.g., Cisco, initially saw very strong demand, but this demand would ultimately not prove sustainable. Platforms (e.g., operating systems, internet browsers etc.) were in an odd position as most of it was open source, and monetization was minimal. Applications were what the end users on the internet actually interacted with. During this period, websites like Google and Amazon were founded.

In the early 2000s, the industry was hit by the dot-com / telecom bubble crash. Every layer of the tech stack was certainly hurt by the crash, but the differences are important. In infrastructure, there are a few important considerations. First, there was significant fiber overcapacity. In the fervor of growth, the network owners built too much supply too early. Additionally, most of the supply was funded via debt. Post crash, unable to meet their obligations, a large number of companies had to ultimately declare bankruptcy. It took nearly a decade after the crash before demand finally caught up to supply and new fiber investments were made. The vendors to the operators (perhaps the closest metaphor to today's chip suppliers) were initially in very high demand as the network operators poured huge investments into the space. They made significant amounts of money during this initial buildout phase as they functioned in an oligopoly and had differentiated offerings, in contrast to the network owners who competed on price. However, when the crash occurred and the fiber investments dried up, the vendors also saw a significant drop in their revenues. It's important to note that the expected internet traffic to meet the supply of internet fiber did ultimately come about and in fact, has well exceeded expectations – but the growth curve was simply slower than forecast.

As mentioned, the platform layer was not really monetized. A lot of the offerings here were open source or bundled with products from other layers. Consequently, most of the investment capital, as well as the subsequent crash, focused on the infrastructure and application layers.

The application layer included the dot-com names that infamously fell in the crash, such as pets.com, that have become synonymous with the irrational exuberance the era was known for. Many of these companies did not survive the crash, and justifiably so, given their lack of sustainable business models and poor unit economics. However, there were several companies that, whilst they did suffer short-term valuation issues, ultimately went on to be multi-decade compounders. Those compounders tended to be the ones with a durable business model, strong unit economics (even if profitability was not yet fully realized), ownership of a critical workflow, and a decent balance sheet.

Ultimately, the infrastructure providers bore the brunt of the capital intensity and the debt, effectively subsidizing the build-out of the internet. While they laid the tracks, the massive economic returns accrued to the applications, aka the Googles and Amazons, that ran on top of them.

Transition to the Cloud

Cloud is the next big tech cycle – one that continues through present day. All three layers for cloud are being built out as services (IaaS, PaaS, SaaS). Infrastructure is focused on building out computing services, offering it as an on-demand utility without any requirement for physical hardware. The physical burden is placed on providers like Amazon Web Services, Microsoft Azure or Google Cloud, who are responsible for the capital investment and maintenance but benefit from immense scale. These are obviously essential parts of the stack, and while the competitive environment is better than it was with fiber, there's enough competition that we'd expect pricing power to decline over time.

Platforms serve as an in-between for applications and infrastructure. While infrastructure provides raw computing power and applications focus on differentiating software for end users, there is a need for a platform in between that helps application software use that power in standardized ways (e.g., scaling up and down, security, and database management; companies include Cloudflare, Snowflake, Datadog etc.). Platforms help application developers focus solely on the actual application software, allowing for easier and faster deployment, testing and launch. These players are important, particularly with application growth, but because they do not actually own the end customer, we believe they are likely to be commoditized over time or bundled with infrastructure. There are similarities in their role in the stack with operating systems in the earlier PC era.

The application layer is software and what most users will be familiar and directly interact with. Software can be broken down into B2B or B2C, sold to companies or to individuals directly, respectively. It can also be broken down into horizontal software, applying across industries (e.g., NOW, WDAY) or vertical software, specific to one industry (e.g., GWRE – P&C insurance, VEEV – healthcare). This has been the layer where a lot of value has accrued over the last 1-2 decades and we’ve seen multiple strong compounders. The best players in this area have been mission-critical systems of record. For these high-quality companies, unit economics are strong and margins are or have the potential to be high. It’s important to note that while we expect many of the established leaders to naturally extend their reach into the AI tech stack, we expect others, particularly those lacking deep moats and mission-critical qualities, to be easily displaced.

The AI Stack

Big Picture – The AI Stack of Today & Tomorrow

As mentioned earlier in this letter, the AI stack is composed of compute and processing power at the infrastructure layer, the harnessing of that power into intelligence at the platform layer, and applications turning the intelligence into real world outcomes at the application layer.

Over the last few years, we’ve seen most of the focus, both in terms of monetary investments and media noise, targeted on the infrastructure layer. We feel this follows similar patterns of past tech cycles – where the infrastructure layer sees significant early benefit as we build out the industry. Demand today clearly outstrips supply, as can be seen in NVIDIA’s incredibly strong growth and margins. However, following the arc of historical cycles, we would expect to see growth moderate mid- to late-cycle due to several reasons. While NVIDIA enjoys a technological edge today, it seems likely that competitors will eventually reach performance sufficient to compete and add pricing pressure. One potential counter to this point is CUDA, NVIDIA’s platform that allows developers to code their GPUs for non-graphic purposes. This would potentially make NVIDIA’s chips much stickier, and it could be very difficult for models to switch to a different company’s chips, even if comparable in features and performance. Increased efficiencies at the platform level will also moderate demand for chips. In fact, we may have already seen early indications of this with DeepSeek.

Quick side note: DeepSeek, for those unfamiliar, is a Chinese AI company that developed its own LLM. Using significantly less money and prior-generation chips (in part due to the export restrictions to China), its LLM results appear competitive to other leading LLMs. It did so by using more efficient training methodologies and cost optimization. It has also made many of its models publicly available, allowing developers to use and build

on it (versus the more closed nature of other leading LLMs). It's perhaps an interesting point of comparison to Linux, which is an open-source operating system. Linux was initially released in the early 90s and was widespread by the end of the decade. As mentioned earlier, the platform layer during the Internet era was largely open and not monetized, with Linux a major driver of this shift.

The platform layer may evolve in a similar way to how it always has – important but fairly commoditized. With several mega tech companies pouring billions or even trillions of dollars into the space (however notably all based on the same architecture), we would expect several competitive models (perhaps even including a few smaller ones who are able to operate more efficiently along the lines of a DeepSeek) that could arguably be interchangeable as the underlying model to an AI application.

Based on the eventual moderation and relative commoditization of the infrastructure and platform layers, we view the application layer as where the long-term value will accrue. This layer will see the rise of new companies both with new and existing use cases that are AI-native, as well as the evolution of existing cloud companies that have successfully incorporated AI. In the best hands, AI is technology that allows an existing company to further their product offering while also being a potential accelerant to internal cost cutting and efficiency gains. To be clear, there will also be many companies in the existing application layer that will fail. At present, we believe these failures are more likely to be the best of breed, one-off, or nice to have offerings that are considered “non-core.” What will remain sticky, assuming they embrace AI, will be the systems of record, suite, and mission critical applications.

The Details – Infrastructure & Platform Spending

We've alluded to similarities between current AI spending and the fiber optic network buildout of the early 2000s. We're focused on two main points: supply vs. demand and funding.

As mentioned, in the late 90s, there was a huge influx of capex into fiber to build out as much supply as possible. This led to significant overcapacity, as a) demand took longer than expected, and b) technology and efficiencies improved so that a strand of fiber could take on actually much more traffic than before. For several years post the crash, as much as 80-90% of the fiber built out during this time remained unused (aka dark fiber).

We are currently seeing huge amounts of money being poured into AI related infrastructure. The question is obviously whether or not we're seeing a repeat of history or if it will be different this time. There are undoubtedly some tech efficiencies that we will see, both in terms of chip usage (as mentioned previously,

DeepSeek as an example) and processing power / data centers, but we think the primary determinant on possible overcapacity will be around demand.

On demand, the internet today has well exceeded long-term expectations and brought possibilities that were probably unfathomable in the late 1990s, but initial traffic took longer than initially forecasted. In the meanwhile, network operators and their vendors suffered greatly due to this mismatch in supply and demand. There are several possibilities as to why demand was slower than expected, ranging from the expectations were just wildly inflated to slow human adoption / behavior to lack of content that demanded significant bandwidth (e.g., video streaming was non-existent / limited).

When considering AI and its capacity to meet the significant infrastructure supply being created, we see a similar path to the internet. Its long-term possibilities are probably beyond what we can appreciate today, but we also think adoption will likely be slower than expected, due to friction in adoption and the reality of AI's capabilities today.

Human behavior is slow to change, and structurally, our societies are slow to do it – particularly when facing potentially revolutionary change. Examples include the internet as noted, on-premise to cloud took / is taking years if not decades, mobile commerce has been a gradual behavioral shift over the last 3 decades etc. Change always faces an uphill battle against the status quo, particularly for institutions, requiring proof of superiority or necessity, establishment of trust, and re-learning of behavior or workflows. These are obviously all well-known considerations that are factored into forecasts, but we also believe that human behavior tends to favor optimism and a fear of underinvesting or missing out outweighs the reverse. As it relates to AI, companies appear to view, probably accurately in some ways, underinvesting as potentially an existential threat while overinvesting could be painful monetarily but likely survivable (especially for cash flow positive mega caps).

On AI's capabilities, while impressive and dramatically improving daily, we still believe we are early in its possibilities, adoption and real-world outcomes, particularly at the enterprise level. While there are often one-off headlines about AI tools allowing a lay person to build a functional software tool over the weekend where previously it would take multiple engineers months, we balance that out with reports like MIT's 2025 research that suggest the significant majority of early genAI projects it reviewed failed to deliver measurable return to their organizations (with significant caveats, refer to The GenAI Divide: State of AI in Business 2025). There's a significant difference between an AI tool being able to output a deliverable correctly to actually using it to replace a company's workflow and being trusted to reliably do so. We also believe it's important to note here again that AI models are probabilistic systems rather than deterministic rule-based software – this allows for

edge cases and situations where models can behave confidently when they should not. The acceptable error threshold obviously varies wildly by industry, but we also believe AI is held to a meaningfully higher standard vs. humans. Said another way, human error is viewed as part of human nature, whereas AI errors can quickly undermine trust and slow adoption. Anecdotally, our own experience with consumer-facing LLMs has supported this view, but we've also spoken to industry experts who have expressed hesitancy around the realistic capabilities of AI today and the gap between one-off demonstrations and reliable enterprise-scale deployment.

A quick counter is that more of AI's demand is built in or at the very least can be pushed by the infrastructure builders. Whereas for the internet, traffic depended on consumers and other enterprises to adopt new behavior. Many of the key players building out AI infrastructure will also be key buyers of its capabilities. Having said that, we do feel strongly the end consumer or enterprises will ultimately still need to build that demand. The increasing efficiency in chips or processing power mentioned earlier, while potentially decreasing the need for infrastructure capacity, can also mean increasing or accelerating potential use cases as lowered costs increases accessibility.

Switching to funding, the fiber optic buildout was funded largely via debt that the network operators took on, as well as equity issuance and vendor financing between equipment suppliers and the network operators. When the crash occurred, many operators were over-levered and had to declare bankruptcy or restructure.

At first glance, funding for AI today looks very different. A significant amount of the capex is being funded by internally generated cash flows from mega-cap tech companies who have meaningful profitability from other lines of business (e.g., Meta, Google etc.). There's also substantial private funding (e.g., VC, sovereign wealth funds etc.). Leverage across the ecosystem appears materially lower than during the telecom buildout. However, there are emerging signs of vendor financing and potentially circular funding dynamics that echo the relationships between equipment vendors and network operators during the fiber optic era. Broadly, our take is that funding is relatively healthier this time around, but we remain cautious – particularly given concentration risk. A handful of key parties currently control the majority of both the spend and funding.

The Details – Addressing the Death of Software

The media has often proclaimed the death of software companies at the hands of AI. We'd like to delve in further and address some of the specific points.

An often-touted idea is that AI allows software code to be written much faster, and consequently, we would see the rise of many new companies that could compete with existing players, potentially slowing revenue growth of existing top players and compressing margins. While we agree that AI allows new software to be coded much faster, we don't think this means a dramatic shift in the competitive landscape. For one, existing players can also use AI and roll out additional offerings with speed, extending their moat. Additionally, beyond writing just the code that replicates a specific use case, software needs to be compatible and integrate with the rest of the system, along with meeting security and compliance standards. It's not just the code itself; it's also data gravity, operational and maintenance experience, established relationships and liability considerations. We'd also note that the software companies most vulnerable to new competitors are the same types of software that have always been most at risk, aka non-critical, non-core, one-off tools, with minimal workflow or data ownership. Meanwhile, we think that large, ingrained software that is deeply embedded in a company's critical operations will continue to remain sticky (even if slightly decreased) and difficult to both replicate and replace.

However, this brings us to the next consideration – does AI make switching software easier? Switching costs can be broken down into a few categories. There are technical issues (integrations, data access etc.), operational friction (workflows, internal reporting systems etc.), cost factors, behavioral considerations (user adaptation) and risk (CYA). AI can probably make it easier to switch from a technical perspective over time, but we believe the other frictions, which arguably are more significant, will remain. As mentioned, we believe the type of software that historically has always been sticky will continue to be sticky, e.g., large enterprise software that operates as a mission-critical system of record and is deeply embedded into a company's key processes. These types of software have large amounts of money invested into them for implementation, customization, employees trained specifically on them etc. That investment needs time to amortize. Reconfiguring a company's financial reporting system, payroll workflows or other internal systems is incredibly difficult and time consuming. Human behavior is slow. Getting an end user retrained and productive on new software takes time. And finally, new software means taking on risk; a major switch can always have unforeseen difficulties and costs, a new software vendor is relatively untested vs. an industry leader, and decision makers are also human – there's no incentive to test a new vendor, absent a substantial price or features delta. Certainly, AI will make it even easier to switch already easy to switch software, but we believe sticky embedded software will remain sticky even if relatively decreased from a technical perspective.

Skeptics of software also point to the current pricing model of many SaaS companies: by headcount. Under this narrative, AI will reduce headcount needs and consequently, SaaS revenue will dramatically decline. We think

this oversimplifies the reality of the workforce and how it has historically evolved to new tech. While overall productivity or output per worker should see great gains, we're not confident that the workforce will simply decline, but rather, we believe an evolution of the types of jobs and responsibilities will occur instead. Regardless, we do not view this as an existential issue but rather a pricing model shift. Even if headcount declines, the value being provided should be substantially more with AI tech incorporated into offerings. This could manifest in either revenue per headcount (offsetting the decline in headcount) or a different type of pricing model entirely. Similar to the on-premises to cloud revenue and margin profile change (e.g., greater revenue and lower margins but net higher profit), we would expect this to be a short-term pricing model shift, but the long-term revenue and/or profit pool would be greater.

A quick note, we've of course seen the recent headlines of companies who've slashed headcount allegedly due to AI, and while those assertions are possibly partially true, it's difficult to separate out actual productivity gains or replacement due to AI vs. broader cost-cutting, cyclical right-sizing that likely would have occurred anyways.

The Reality of the Markets & Where This Could Go Wrong

The past few years, infrastructure names have led the public markets. Nvidia has multiplied many times over. Most AI models are either private or within large tech conglomerates, and the majority of investment in this theme has flowed to the hardware companies like Nvidia. Meanwhile, application software, due to many of the fears outlined above, has dramatically underperformed. Over the course of this period of underperformance we have invested heavily across many core, mission critical application software businesses that we believe the market has been overly punitive to.

We do, however, recognize there are several ways our industry thesis could go wrong. We've alluded to and addressed many of them in the prior paragraphs, but here is a quick summary of a very bearish state of the world for application software: the tech stack works differently in this cycle. Infrastructure, probably in part due to Nvidia's CUDA platform, remains a partial bottleneck, and as such, demand sustains as well as significant margins. It operates as a monopoly or in some sort of strong oligopoly. There are a handful of key models who may or may not be differentiated, but they allow software code to be created and switched fairly interchangeably. As such, software is a highly competitive space, primarily competing in price. Value in this stack accrues mostly at the model or platform layer while chips and infrastructure remain a steady grower due to continued demand. The realistic range of outcomes probably incorporates some of these factors but to varying

degrees (e.g., does a niche productivity tool lose its differentiating power, or frankly become bundled as an add-on feature to a larger application? Probably).

A more existential bearish scenario is that we have fundamentally underappreciated the speed at which machine intelligence could replace human intelligence. In this case, the workforce is cut across industries far beyond expectations while overall productivity gains and output dramatically increase, driven by AI. If workflows reach a point where no or minimal human intervention is required, this would substantially destroy one of software's key moats, as there would be minimal to no switching costs. Our concept of "work," as we know it today, could look completely different. However, history suggests that humans do not simply surrender their agency to new tools; we use them to expand the frontier of what is possible. From the steam engine to the internet, tech revolutions haven't just automated old tasks. They have allowed us to push previously held boundaries, creating new tasks and possibilities that require more human oversight and ingenuity, not less. We find the total displacement of human intelligence, or even anything close to it, to be highly unlikely, particularly within our investment horizon.

There's also the possibility our thesis is right, but we've timed our investments poorly. As long-term investors, we generally speak in hypotheticals about timing. Our viewpoint is to think about the long-term outcome or evolution of an industry. Whether it reaches that state in 7 years or 10 years or 15 years is theoretically less consequential, and for the most part, we are willing to look wrong in the short term to be right in the long term. Having said that, should we or could we have predicted that in the early stages of the AI cycle, which we are likely still in, demand and revenue would be focused on infrastructure? Upon reflection yes, but everything of course seems obvious in hindsight.

While we believe that the application layer will see several long-term compounders through this cycle, it could take years or even decades for those returns to be realized. If we are ultimately proven right, if it takes a decade of dead money, are these still considered good investments? These are considerations we weigh every day, along with continued monitoring and diligence on the fundamentals of our companies, as well as that of the broader industry.

How Has Our Industry Thesis Manifested into Investments

To restate, our core industry thesis is that true long-term compounders will emerge at the application layer, where companies own the mission-critical workflows, proprietary data, and customer relationship. The majority of our investment focus in this space has been on identifying those companies.

In horizontal, we are currently invested in Workday and ServiceNow, offering HR / financials and workflow software respectively. Workday is a system of record for employee and financial data. ServiceNow was a system of action initially but has also become a system of record for operational data. Interestingly, this is a name that is often touted as most at risk of AI displacement. However, we think AI will actually be a driver of NOW's value. AI should digitize more workflows out of human hands, and NOW is the orchestration layer that can do that, whether or not those tools are from NOW itself. NOW is the tech infrastructure in a company that can connect various tools together and piece together appropriate workflows (the risk of AI is really to those individual standalone tools).

Vertical software is arguably even better positioned to withstand threats, given the deep industry expertise and specialization required. We are currently invested in Guidewire and Veeva. Guidewire's platform manages policy, billing and claims for P&C insurance companies. Veeva is the dominating CRM platform in the healthcare space and has a growing business to manage R&D lifecycle. Both operate in highly regulated industries with significant compliance and security requirements. This is not a space where we would expect untested vendors to make significant strides quickly.

All of our names are incredibly sticky, with churn rates in the low single digits. They are also all incorporating AI into their existing offerings, whether through M&A or internal investments.

It's difficult (or frankly impossible) to try to forecast what the exact state of the world will be in a decade or two. However, we feel confident with the diligence we have done that the current companies we are invested in can successfully navigate and compound profitably in a large range of realistic, possible outcomes. As always, we continue to test our assumptions, challenge our views, and actively monitor developments.

We look forward to continuing to update you on our portfolio and progress.

Warm Regards,

Lucy and Nate

06/01/2026

Appendix – Performance Since Inception

	Q4 2021	CY 2022	CY 2023	CY 2024	CY 2025	Q1 2026	Since Inception
Luna Partners (Net)	5.6%	-24.6%	36.7%	-3.4%	12.6%	-24.6%	-10.7%
Luna Partners (Gross)	5.8%	-23.7%	38.7%	-1.8%	14.7%	-24.2%	-4.4%
S&P 500	10.6%	-19.4%	24.2%	23.3%	16.4%	-4.6%	51.6%
S&P MidCap 400	7.6%	-14.5%	14.4%	12.2%	5.9%	2.2%	27.9%
S&P SmallCap 600	5.3%	-17.4%	13.9%	6.8%	4.2%	3.1%	13.6%
Russell 3000	8.9%	-20.5%	24.0%	22.1%	15.7%	-4.3%	45.2%

Note: Equity market returns above based on underlying index prices. Index prices are sourced from Koyfin, Yahoo Finance or other publicly available sources. S&P 500 is a popular stock market index representing 500 large companies. S&P MidCap 400 is a stock market index representing 400 mid-sized companies. S&P Small Cap 600 is a stock market index representing 600 small-sized companies. The S&P indices are composed of companies on the US stock exchanges, and each S&P index has distinct constituents. Russell 3000 is a stock market index representing the 3000 largest companies on the US stock exchanges.

Luna Partners' inception date is October 1, 2021. Luna's net performance is based on net asset value calculations provided by Interactive Brokers and calculated net of expenses and fees. Luna's gross performance is calculated before fees to Luna. Luna's portfolio typically consists of 10-15 long, US based public equity positions. Returns are presented on a time-weighted basis to be comparable to the index returns in the table. Past performance is not indicative of future returns.

Luna Partnership LLC is a state registered investment adviser located in Miami, Florida.

The views outlined in this newsletter are those of Luna Partners and should not be construed as individualized or personalized investment advice. Any economic and/or performance information cited is historical and not indicative of future results. Economic forecasts set forth may not develop as predicted. Although the information in this report has been compiled from data considered to be reliable, the information is unaudited and is not independently verified. We do not guarantee that the forgoing material is accurate, complete, or up to date. As a result, there's a risk that the information is not accurate.

Performance results reflect the deduction of advisory fees, brokerage commissions, and fund charges. Performance results reflect the investment of dividends and other earnings. This information is provided to you in a combined form, solely for your convenience and ease of review. Past performance is never a guarantee of future results. All investments have investment risks such as fluctuation in investment principal including the complete loss of principal invested.

Different types of investments involve varying degrees of risk, and there can be no assurance that the future performance of any specific investment, investment strategy, or product made reference to directly or indirectly, will be profitable, equal any corresponding indicated historical performance level(s), or be suitable for a given client or portfolio. Investing in stock includes numerous specific risks including: the fluctuation of dividend, loss of entire principal and potential illiquidity of the investment in a falling market.

Any questions regarding the applicability of any specific issue discussed above should be addressed with Luna Partners. All information, including that used to compile charts and/or tables, is obtained from sources believed to be reliable, but Luna Partners has not verified its accuracy and does not guarantee its reliability. Different types of investments involve varying degrees of risk, and there can be no assurance that the future performance of any specific investment, investment strategy, or product made reference to directly or indirectly in this newsletter, will be profitable, equal any corresponding indicated historical performance level(s), or be suitable for your portfolio.

Moreover, you should not assume that any discussion or information contained in this newsletter serves as the receipt of, or as a substitute for, personalized investment advice from Luna Partners or from any other investment professional. To the extent that you have any questions regarding the applicability of any specific issue discussed above to your individual situation, you are encouraged to consult with Luna Partners or the professional advisor of your choosing. All information, including that used to compile charts, is obtained from sources believed to be reliable, but Luna Partners has not verified its accuracy and does not guarantee its reliability.

This newsletter does not constitute an engagement with Luna Partners. The information contained in this newsletter is general in nature and offered only for educational purposes. No investment decisions should be based upon this newsletter's contents. This newsletter is not a substitute for engaging Luna Partners for a personal consultation whereby all of your financial circumstances, risk tolerance and investment objectives can be considered. Information pertaining our investment advisory operations, services, and fees is set forth in the Advisor's Form ADV Disclosure Brochure, a copy of which is available from your investment adviser representative upon request. Please contact us at 201-953-5874 or by email at Ichou@luna.partners with any questions you may have, or if you would like to review this report or your current investment strategy.

This newsletter was created by Luna Partners utilizing information provided by Interactive Brokers. The pricing for the investment positions listed in the report is based upon information received from the account custodian. The performance listed in this report may vary from the information provided on your account statement provided by your account custodian. Luna Partners is independent of Interactive Brokers. Luna Partners is relying upon the information provided by Interactive Brokers, and Luna Partners has not audited or otherwise verified the accuracy of the methodology, calculations or information in this report. As a result, the methodology, calculations and information presented in the report are not guaranteed by Luna Partners.