

Flux Core Data Systems and Energy

Table of Contents

1. Offtake Agreements: Turning Megawatts into Bankable Demand
2. Bare Metal Provisioning: The Powered Shell Rebuilt for AI Density
3. Power as a Service: Buying Power Without Owning the Bottleneck
4. Oil and Gas Well Deployments: Turning Flared Gas into AI Compute
5. Zero Carbon Footprint: Engineering Reductions, Not Buying Offsets
6. Sovereign AI, Sovereign Data: Compute That Stays Inside the Perimeter
7. College and University Data: Research Compute Without the Grid Wait
8. Powered Land Lease Rates: The Real Bubble in AI Infrastructure

1. Offtake Agreements: Turning Megawatts into Bankable Demand

Executive Summary

The bottleneck in artificial intelligence infrastructure is no longer whether megawatts can be built. It is whether those megawatts are spoken for before the first breaker closes. A data center without contracted demand is a speculative real estate bet, and capital markets treat it accordingly. Flux Core Data Systems and Energy leads every deployment with a bankable offtake position. This paper explains why contracted demand has become the true unit of value in AI infrastructure, why buyers increasingly cannot secure guaranteed capacity in a supply constrained market, and how Flux Core converts megawatts into financeable, revenue backed commitments. We continue to seek additional offtake partners across neoclouds, enterprises, universities, hospitals, municipalities, national labs, and AI platforms.

The Problem

Demand for compute is exploding while the ability to serve it lags badly. Goldman Sachs estimates that United States data center power demand will rise from roughly 31 gigawatts in 2025 to about 66 gigawatts by 2027, more than doubling in two years (Source 1). In PJM, the largest grid operator in the country, data centers are expected to account for 30 of the next 32 gigawatts of load growth, a concentration that reveals just how dominant AI has become as a driver of new electricity demand (Source 1). When a single sector consumes almost all incremental grid capacity in the nation largest market, the ordinary rules of supply and demand no longer apply. Capacity stops being something a buyer can procure on short notice and becomes something a buyer must reserve years in advance or forfeit entirely.

That scarcity changes what a lender or investor is actually underwriting. In a balanced market, a data center is a real estate asset whose value derives from a predictable stream of tenants. In a supply constrained market, the asset value collapses onto a single question: is the capacity contracted? A speculative build, however well engineered, carries the full risk that the developer never signs paying customers, that the interconnection slips, and that the equipment arrives before the demand does. Capital providers price that risk punitively, or they decline to fund the project at all. The offtake agreement is what converts a speculative

build into a financeable one, because it substitutes a contractual revenue stream for a market bet.

Yet the constraint is not generation capacity in the abstract. It is the grid itself. Analysis from S&P Global Market Intelligence concludes that transmission and interconnection, not power plants, are the true bottleneck for new data center capacity, and that United States data center capacity will need to grow from about 62,242 megawatts in March 2026 toward roughly 151,734 megawatts by 2030 to keep pace (Source 2). When supply is that constrained, the scarce commodity is not a building. It is deliverable, contracted capacity that a buyer can actually count on.

This creates a structural mismatch. Developers cannot finance projects without committed demand, because lenders and investors will not underwrite speculative megawatts in a market where interconnection timelines and equipment lead times routinely stretch past the life of a typical financing assumption. At the same time, buyers of compute, from neoclouds reselling GPU hours to enterprises training private models, cannot secure guaranteed capacity because everyone is competing for the same scarce, deliverable megawatts. Capacity without committed demand is unfinanceable, and demand without deliverable capacity is unfulfillable. The offtake agreement is the instrument that resolves both sides of that mismatch at once.

Real-World Example

The public data makes the tension concrete. With United States data center power demand projected to more than double to 66 gigawatts by 2027 (Source 1) and grid interconnection identified as the binding constraint (Source 2), a buyer that waits for capacity to appear on the open market is likely to be disappointed. The winners are the parties that lock capacity in advance. S&P Global projects that United States data center capacity must grow from about 62,242 megawatts in early 2026 to roughly 151,734 megawatts by 2030, a near tripling that the grid cannot deliver on its own timeline (Source 2). Every megawatt of that growth is a megawatt that someone will need to have contracted ahead of time.

Consider an illustrative scenario. A regional neocloud has signed inference customers and a growing training pipeline, but no way to guarantee the GPU capacity its own customers are demanding. Going to the public grid means entering an interconnection queue with a multi year wait, an outcome that would force the neocloud to turn away revenue today.

Instead, the neocloud reserves capacity inside a Flux Core deployment, securing a defined block of megawatts with a delivery timeline measured in months rather than years. That reservation becomes a demand signal that anchors the site. This scenario is illustrative and does not describe a specific Flux Core customer, but it mirrors the exact behavior the market data predicts.

How Flux Core Solves It

Flux Core makes megawatts bankable by making them deliverable and diversified. Our demand book is deliberately spread so that no single sector dictates the timeline or the risk profile. Neoclouds reselling GPU hours, enterprises running private models, universities and national labs, hospitals, municipalities, and AI platforms scaling inference each represent an independent stream of contracted demand. A diversified offtake book behaves like a diversified revenue portfolio, and it is far more financeable than a single tenant bet. If one sector softens, the others carry the deployment, which is exactly the resilience that capital markets reward with better terms.

This diversification is also a strategic hedge against the concentration risk the market is building into itself. When PJM reports that data centers account for 30 of the next 32 gigawatts of load growth (Source 1), it signals a market where a handful of hyperscale buyers dominate demand. A deployment anchored to a single such buyer inherits that buyer's entire credit and strategy risk. Flux Core deliberately builds a broader base, including institutions such as universities, hospitals, and municipalities whose demand is durable and whose planning horizons are long, so that contracted capacity rests on many shoulders rather than one.

Deliverability is the other half of the equation, and it rests on two engineering choices. First, Flux Core deploys the Nexus R1000, a 1MW containerized data center that comes online in three to six months rather than the twelve to twenty-four month timelines of conventional construction. Reserved capacity that arrives in months is a fundamentally different product from capacity that arrives in years. Second, Flux Core is power agnostic. Reserved capacity can sit on stranded or associated natural gas, on renewable solar paired with battery energy storage, on a hybrid microgrid, or on an existing grid tie. A commitment therefore never stalls behind a multi-year interconnection request, which is precisely the risk that S&P Global identifies as the binding constraint on new capacity (Source 2).

Underpinning all of it is a closed loop liquid cooling system that uses zero outside water. Because the fluid recirculates with no cooling tower and no evaporative loss, there is no make up water permit to secure and no dependence on local water supply. That removes an entire category of permitting and siting risk from the delivery path, which is exactly what a lender or an anchor tenant needs to see before signing. The Nexus platform is built to Tier III design standards, targeting 99.982 percent uptime, so the contracted capacity is not only deliverable but dependable.

Put together, the Flux Core model reframes what an offtake partner is actually buying. It is not square footage and a hope that power arrives. It is a defined block of capacity with a defined delivery date, an energy path that does not depend on a queue, and a cooling architecture that removes the water and permitting risks that so often derail conventional projects. For a lender, that combination is what turns a contracted megawatt into collateral. For a buyer, it is what turns a reservation into a guarantee.

Conclusion and Call to Action

In a market where United States data center power demand is set to double by 2027 (Source 1) and the grid is the bottleneck (Source 2), the offtake agreement is the instrument that turns scarce megawatts into bankable demand. Flux Core builds every deployment around that instrument, backed by containerized delivery in months, a power agnostic energy path, and zero outside water cooling. If you need guaranteed AI capacity, or if you want to anchor a site as an offtake partner, reach out to Flux Core to see what megawatts remain open. Own the Energy, Monetize the Data.

Sources

1. TheStreet, AI energy appetite and the electric grid (Goldman Sachs figures).
<https://www.thestreet.com/markets/ais-energy-appetite-is-reshaping-the-electric-grid-data-centers>
2. S&P Global Market Intelligence, AI data center power demand and grid constraints.
<https://www.spglobal.com/market-intelligence/en/news-insights/research/2026/06/ai-data-center-power-demand-grid-constraints-energy-resilience>

2. Bare Metal Provisioning: The Powered Shell Rebuilt for AI Density

Executive Summary

The powered shell was a good idea for a different era. It gave tenants power and walls and let them finish the interior on their own schedule. That model breaks completely under the weight of modern AI hardware, where a single rack can draw more power than an entire row did a decade ago. Flux Core Data Systems and Energy has rebuilt the powered shell for AI density: a bare metal environment provisioned for up to 1.3MW of IT load, engineered for the hottest and densest GPU workloads, cooled by a closed loop liquid system that consumes zero outside water, and deployable in three to six months. This paper explains why conventional shells fail at AI density, quantifies the timeline and lead time problems, and details how the Flux Core approach delivers dense, deploy ready white space on a timeline the market cannot otherwise match.

The Problem

Rack density has changed by an order of magnitude and then some. Industry observers note that rack power density has risen roughly fiftyfold over the past decade, driven by accelerated computing and the shift to GPU based training and inference (Source 1). A conventional powered shell was designed for cabinets pulling a handful of kilowatts. The current generation of AI systems, such as GB300 NVL72 class deployments, can demand 130 to 140 kilowatts per rack, a load that air cooled halls and standard shells simply cannot dissipate.

The second failure is time. A conventional powered shell typically takes twelve to twenty four months to go from a signed lease to a rack a tenant can energize, and that assumes power is available. In practice, it often is not. Analysis of interconnection queues finds that PJM interconnection has averaged around 40 months, with data center heavy zones running 36 to 48 months, even as federal targets aim to compress that to a fraction of the time (Source 2). Grid and equipment lead times compound the delay: turbine and transformer lead times now run several years, and gas nameplate additions remain limited relative to demand (Source 3). A tenant with GPUs on order and customers waiting cannot absorb a two to four year power path layered on top of a two year build.

These two failures, density and time, reinforce each other in a way that makes the conventional shell a dead end for AI. Even if a tenant were willing to wait years for power, the shell that eventually energizes still cannot cool the racks the tenant needs. And even if a shell could somehow be retrofitted for density, the power to feed those denser racks is still years away in the queue. The tenant is trapped between a cooling ceiling it cannot raise and a power timeline it cannot compress. Solving one without the other accomplishes nothing, which is why incremental improvements to the traditional model do not work. The shell has to be reconceived from the ground up around both constraints at once.

The result is a category of stranded demand. Tenants have hardware, capital, and customers, but no dense, powered, deploy ready white space to put the hardware into. The shell they can lease cannot handle the racks they need to install, and the timeline to fix that stretches beyond the commercial window in which the compute is valuable. The conventional colocation industry, built around a slower and less dense era, is structurally unprepared to close that gap on the timeline the market now requires.

The cooling gap deserves particular emphasis, because it is the hidden reason conventional shells fail. Air cooling can reasonably serve racks in the range of ten to twenty kilowatts. Beyond that, moving enough air to remove the heat becomes physically impractical and energy wasteful. Modern accelerators blow past that ceiling immediately. An H200 or B200 class deployment runs at roughly 40 to 60 kilowatts per rack, and a GB300 NVL72 class rack reaches 130 to 140 kilowatts. A powered shell that was designed and permitted around evaporative cooling towers and air handling cannot be retrofitted to those densities without effectively rebuilding it, which reintroduces the very timeline the tenant was trying to avoid.

Real-World Example

The public numbers frame the gap precisely. With rack density up roughly fiftyfold in a decade (Source 1) and interconnection driven power timelines running 36 to 48 months in data center zones (Source 2), a tenant that needs dense racks fast has almost no conventional option. Even where a shell exists, the cooling and power provisioning were designed for a load profile that no longer applies.

Consider an illustrative scenario. A GPU tenant has secured an allocation of high density accelerators and signed customers who expect capacity this year. The tenant evaluates a conventional powered shell and discovers two disqualifying facts: the hall cannot cool racks

above a fraction of the density the hardware requires, and the power path behind the shell is still years from energization. This scenario is illustrative and does not describe a specific Flux Core customer. It reflects the structural mismatch the data describes, and it is exactly the situation the Flux Core powered shell is built to resolve.

The financial stakes of that mismatch are severe. High end accelerators depreciate rapidly, and the customer contracts that justify their purchase have hard start dates. A tenant that cannot energize dense racks on time is paying for idle hardware while its competitors serve the market. The public data on interconnection delays of 36 to 48 months (Source 2) and multi year equipment lead times (Source 3) means the conventional path does not merely cost money. It forfeits the entire commercial opportunity, because by the time the shell is ready, the hardware generation has moved on.

How Flux Core Solves It

Flux Core delivers a bare metal environment provisioned for up to 1.3MW of IT load, pre engineered for the densest AI and GPU workloads. Density is a cooling problem before it is a power problem, and Flux Core solves cooling at the foundation. The closed loop liquid cooling system recirculates a single fluid and consumes zero outside water, which eliminates cooling towers, make up water permits, and the evaporative losses that cap traditional halls. That thermal headroom is what makes extreme density possible. The platform supports Configuration 2 for H200 and B200 class systems at roughly 40 to 60 kilowatts per rack, and Configuration 3 for GB300 NVL72 class systems at roughly 130 to 140 kilowatts per rack. Cabinets run hot chip dense at Tier III design reliability, targeting 99.982 percent uptime.

The shell is also power agnostic, which is what breaks the timeline problem. Whether a site offers stranded or associated natural gas, renewable solar with battery energy storage, a hybrid microgrid, or an existing grid tie, Flux Core builds the matching power path rather than waiting in an interconnection queue. That is how a 1.3MW environment reaches edge and constrained sites that conventional colocation avoids, and it is why deployment happens in three to six months instead of the multi year timelines that grid dependence imposes (Source 2, Source 3).

The container form factor ties it together. Because the Nexus R1000 is containerized, deployable, and relocatable, a tenant is not committing capital to a fixed building on a fifteen

year horizon against a market that changes every quarter. Tenants drop in their own hardware or lease Flux Core hardware, energize dense racks in months, and retain the option to relocate as needs shift.

It is worth stating plainly what the rebuilt powered shell replaces. The old model asked a tenant to accept a slow, capital heavy, grid dependent building that could not cool modern hardware, and to hope that power and cooling could be sorted out later. The Flux Core model delivers cooling and power provisioning as first class design elements rather than afterthoughts. Density is engineered in through liquid cooling, power certainty is engineered in through the power agnostic architecture, and speed is engineered in through the containerized form factor. A tenant no longer has to choose between fast, dense, and powered. The rebuilt shell offers all three at once, which is the standard AI density now demands and conventional colocation cannot meet.

The economic implications for a tenant follow directly from those engineering choices. Because the environment is provisioned for up to 1.3MW of IT load and cooled for the full range of dense configurations, a tenant can pack far more compute into a given footprint than a conventional hall allows, which improves the utilization of both space and capital. Because deployment happens in months rather than years, the tenant begins serving customers and generating revenue from expensive accelerators far sooner, before those accelerators depreciate. And because the units are relocatable, the tenant is not locked into a single site through a long real estate commitment if its market or its hardware roadmap changes. Speed, density, and optionality are not merely operational conveniences. They are the financial levers that determine whether an AI deployment earns a return before the technology moves on.

Conclusion and Call to Action

The powered shell has to be rebuilt for a world where racks draw fiftyfold more power than they used to (Source 1) and grid timelines run three to four years (Source 2, Source 3). Flux Core has done exactly that: a bare metal, 1.3MW provisioned shell with closed loop zero water cooling that enables 130 to 140 kilowatt racks, Tier III design, three to six month deployment, and a power agnostic energy path. If you are deploying high density GPUs on a timeline, ask Flux Core about a ready powered shell provisioned up to 1.3MW. Own the Energy, Monetize the Data.

Sources

1. Core Insights via LinkedIn, rack density growth (roughly fiftyfold in a decade). <https://www.thestreet.com/markets/ais-energy-appetite-is-reshaping-the-electric-grid-data-centers>
2. Carbon Direct, analysis of power grid interconnection queues (PJM, ERCOT). <https://www.carbon-direct.com/press/carbon-direct-releases-new-analysis-of-power-grid-interconnection-queues-pjm-ercot>
3. SemiAnalysis, United States grid constraints toward 40 gigawatts behind the meter. <https://newsletter.semianalysis.com/p/us-grid-constraints-towards-40gw>

3. Power as a Service: Buying Power Without Owning the Bottleneck

Executive Summary

The single hardest thing about building AI infrastructure today is not chips or capital. It is getting power connected. Interconnection queues now stretch past five years in the largest markets, and utilities are increasingly telling data center developers to bring their own power. Flux Core Data Systems and Energy answers with Power as a Service: we finance, own, and operate the generation and distribution, and the customer buys reliable power for compute as a clean operating expense. This paper quantifies the interconnection crisis, explains why behind the meter power is moving from a niche workaround to a mainstream strategy, and shows how a power agnostic, zero water Power as a Service model lets customers buy power without owning the bottleneck.

The Problem

The scale of the interconnection backlog is staggering. Roughly 2,600 gigawatts of capacity sits in United States interconnection queues, with a median wait of about five years to connect (Source 1). That is more capacity waiting to connect than the entire installed base of the grid, and the queue moves slowly. For a data center developer, a five year wait is not a delay. It is a project killer, because the commercial window for AI capacity is measured in quarters, not half decades.

Utilities have begun to say the quiet part out loud. Rather than promising timely interconnection, many are effectively telling developers to bring their own power. Regulators are trying to help. Federal efforts at FERC and DOE aim to cut queue times from more than five years to under two years, and to enable up to 40 gigawatts of new capacity by 2030 against a baseline nearer 15 gigawatts (Source 2). Those reforms are welcome, but they are aspirational, and no developer can build a business on a queue reform that has not yet taken effect.

The market has already drawn its own conclusion. SemiAnalysis projects that grid headroom turns negative around 2027, meaning the grid cannot absorb the incremental load, and forecasts more than 40 gigawatts of behind the meter data center capacity in the United States by 2028 (Source 3). S&P Global expects behind the meter generation to meet

roughly 25 percent of new data center demand by 2030 (Source 4). Behind the meter is no longer a fringe tactic. It is becoming the default answer to a grid that cannot keep up.

What that shift really represents is a rethinking of who owns the power problem. For decades, the assumption was that a data center plugs into a utility and the utility handles generation. The AI era has broken that assumption, because the utility can no longer connect load fast enough, and the equipment needed to expand the grid is itself on multi year lead times (Source 3). Developers now face a choice: wait indefinitely for the grid, or take ownership of generation themselves. Taking ownership, however, means becoming an energy company, with all the capital, expertise, and operational burden that entails. Most compute operators neither want nor are equipped for that role. The gap between needing your own power and being able to run a power plant is precisely the gap Power as a Service fills.

Real-World Example

The public data describes a trap. With about 2,600 gigawatts stuck in queues and a five year median wait (Source 1), and grid headroom projected to go negative around 2027 (Source 3), a developer that depends on the utility is likely to miss the market entirely.

Consider an illustrative scenario. A developer holds a strong site and financing, but the interconnection study returns a five year timeline and a large network upgrade cost. Waiting means forfeiting customers who need capacity now. This scenario is illustrative and does not describe a specific Flux Core customer, but it is the exact situation the queue data predicts for thousands of projects. The developer that instead moves behind the meter, generating power on site rather than waiting in line, converts a five year dead end into a project that can energize in months. That is the strategic pivot the whole market is now making, and it is what Power as a Service delivers as a service rather than a capital project.

How Flux Core Solves It

With Power as a Service, Flux Core finances, owns, and operates the generation and distribution, and the customer buys reliable power for compute as a predictable operating expense. There is no capital locked into generators, switchgear, or storage, and no project timeline held hostage by a utility queue. The customer gets power. Flux Core owns the bottleneck.

Because Flux Core is power agnostic, the system underneath the load can be stranded or associated natural gas, solar paired with battery energy storage, a hybrid microgrid, or a grid connection, sized to the customer ramp and the site. That flexibility is what makes the behind the meter strategy practical across geographies rather than only where one resource happens to be abundant. It aligns directly with the market shift that S&P Global and SemiAnalysis describe, where behind the meter moves toward a quarter of new demand and past 40 gigawatts of capacity (Source 3, Source 4).

The operating expense structure matters as much as the technology. A conventional approach forces a developer to raise and deploy capital for generators, switchgear, transformers, and storage, then carry the operating burden and the technical risk of running a power plant, a business most compute operators have no desire to be in. Power as a Service removes that entire distraction. The customer keeps its capital and its focus on compute, while Flux Core takes on the generation, the maintenance, the fuel procurement, and the balancing. Predictable per unit pricing replaces lumpy capital outlay and unpredictable utility rate risk, which is a materially better profile for planning and for financing.

This model also insulates the customer from the reforms that have not yet arrived. The federal effort to compress queue times from five years to under two years (Source 2) is real and worth supporting, but no operator can build a launch plan around a reform that is still working through rulemaking and implementation. Power as a Service lets a customer proceed today on infrastructure Flux Core controls, and then, if and when grid access improves, incorporate a grid tie as one more option within the same power agnostic system. The customer is never stranded waiting for policy to catch up to demand.

Thermal efficiency holds up alongside the power model. Closed loop liquid cooling with zero outside water keeps PUE low and takes the facility off local water supply, a decisive advantage in the arid, low cost power regions where compute wants to be. The Nexus R1000 delivers this in a 1MW containerized package built to Tier III design standards, deployable in three to six months. The customer pays a predictable rate while Flux Core monetizes the underlying energy assets, credits, and tax structures, turning the power bottleneck from a liability into a managed service.

The deeper value of Power as a Service is that it lets each party do what it does best. Compute operators are experts in silicon, software, and serving customers, not in permitting turbines, negotiating gas supply, or balancing a microgrid. Flux Core specializes in exactly those energy disciplines. By separating the power problem from the compute problem and taking ownership of the former, Power as a Service removes the single biggest source of delay and risk from an AI infrastructure project, and it does so without asking the customer to become an energy company. That division of labor is why the model scales, and why it fits customers from neoclouds to hospitals to national labs regardless of their in house energy expertise.

Resilience is an underappreciated benefit of the model. A behind the meter system that Flux Core owns and operates can be engineered for the reliability that mission critical compute requires, with redundancy and Tier III design standards built in, rather than inheriting the fragility of an overtaxed regional grid. When the grid is strained, as it increasingly is in the very regions where compute demand concentrates, a facility that generates its own power is insulated from curtailment, brownouts, and the volatility of wholesale electricity prices. The customer gets not only power that arrives on time, but power that stays on and stays predictable, which is precisely the profile that AI workloads with hard uptime requirements need. That combination of speed, predictability, and resilience is what elevates Power as a Service from a financing convenience to a genuine operational advantage.

Conclusion and Call to Action

When 2,600 gigawatts sit in queues with a five year wait (Source 1) and behind the meter is becoming the default (Source 3, Source 4), the smart move is to stop owning the bottleneck and start buying power as a service. Flux Core finances, owns, and operates power agnostic generation cooled by a zero water closed loop system, so customers energize in months and pay as they go. If you have been told to bring your own power, ask Flux Core how Power as a Service delivers it. Own the Energy, Monetize the Data.

Sources

1. ChargedUp, the 2,600 gigawatt interconnection queue and time to power. <https://chargeduppro.com/post/2600-gigawatt-interconnection-queue-time-to-power-2026>
2. Gentis News, FERC and DOE fast track for AI data centers. <https://gentis.news/article/ferc-doe-fast-track-ai-data-center>
3. SemiAnalysis, United States grid constraints toward 40 gigawatts behind the meter. <https://newsletter.semianalysis.com/p/us-grid-constraints-towards-40gw>

4. S&P Global Market Intelligence, behind the meter to meet about 25 percent of new demand by 2030.
<https://www.spglobal.com/market-intelligence/en/news-insights/research/2026/06/ai-data-center-power-demand-grid-constraints-energy-resilience>

4. Oil and Gas Well Deployments: Turning Flared Gas into AI Compute

Executive Summary

Every year, the oil and gas industry burns off enough natural gas to power a significant share of the world economy, simply because there is no economical way to move it. That wasted energy is a stranded asset and an environmental liability at the same time. Flux Core Data Systems and Energy resolves both by deploying containerized data centers at the wellhead, converting stranded and associated gas into AI compute revenue on site, with no midstream takeaway required. This paper quantifies the scale of global flaring, explains why so much gas is stranded, and details how a wellhead deployment powered by that gas, cooled by a zero water closed loop system, and relocatable as production shifts turns a liability into a new revenue line.

The Problem

Global gas flaring reached 167 billion cubic meters in 2025, the highest level since 2019 and a rise for the third consecutive year, according to the World Bank Global Gas Flaring Tracker as reported by Energy Connects (Source 1). That flared gas represents roughly 54 billion dollars in wasted energy and about 429 million tonnes of carbon dioxide equivalent in emissions (Source 1). The Payne Institute analysis of the same World Bank data confirms the upward trend and the persistent gap between flaring reduction pledges and actual outcomes (Source 2). The United States remains among the top flaring countries in the world.

The reason so much gas is burned is not carelessness. It is economics and infrastructure. When an oil well produces associated gas, that gas has to go somewhere. If there is no pipeline takeaway capacity, or if the gas is too small in volume or too remote to justify a midstream connection, the operator faces a choice between shutting in valuable oil production and flaring the gas. Flaring is often the only legal and practical option. The molecule that could have powered something useful instead goes up the stack, wasting energy and generating emissions with zero economic return.

The geography of the problem compounds it. Stranded gas is, by definition, produced where demand is not. Building the midstream infrastructure to connect that gas to a market can take years and require permits, rights of way, and capital that the volume of a single

pad may never justify. Meanwhile, oil production continues, and the associated gas keeps coming. The operator is left with a resource that has real energy value but no economical path to a buyer. This is why flaring has proven so stubborn despite decades of reduction pledges, and why the Payne Institute finds flaring rising for a third consecutive year (Source 2). The problem is not a lack of will. It is a lack of a way to convert a remote molecule into value at the point of production.

This is a structural mismatch between where energy is produced and where it can be used. Traditional monetization requires moving the gas to demand through expensive, slow to permit midstream infrastructure. In arid producing basins, there is an additional constraint: water. Conventional data centers depend on evaporative cooling and make up water, which is scarce and contested in the very regions where stranded gas is most abundant. Any solution that requires trucking in water or securing a water permit fails before it starts.

The regulatory and reputational pressure on flaring is intensifying at the same time. Investors, lenders, and regulators increasingly scrutinize the emissions intensity of production, and routine flaring is a visible, quantifiable liability that shows up in environmental reporting and, in some jurisdictions, in permitting decisions. The Payne Institute analysis underscores that flaring has now risen for three consecutive years despite widespread reduction commitments, which means the gap between pledges and outcomes is widening rather than closing (Source 2). An operator that can demonstrably reduce its flare while creating a new revenue stream converts a growing liability into an asset, which is a materially stronger position than simply buying offsets or waiting for pipeline capacity that may never come.

Real-World Example

The scale is undeniable. With 167 billion cubic meters flared in 2025, roughly 54 billion dollars wasted, and 429 million tonnes of carbon dioxide equivalent emitted (Source 1, Source 2), the stranded energy problem is enormous and growing.

Consider an illustrative scenario. A Permian operator produces oil from a pad with substantial associated gas, but no pipeline takeaway. The operator watches a year of that gas go up the flare stack because moving it is uneconomical and shutting in oil is unthinkable. That gas could have been running AI compute. This scenario is illustrative and does not describe a specific Flux Core customer, but it reflects a situation repeated across

producing basins worldwide, consistent with the flaring volumes the World Bank reports (Source 1). Bloomberg and other outlets have documented how persistent takeaway constraints keep flaring elevated even when operators would prefer to capture the gas.

Extend the scenario to see the reversal Flux Core enables. The same operator, instead of flaring, hosts a containerized deployment on the pad that consumes the associated gas as fuel for on site generation. The gas that had negative value now powers revenue generating compute. The flare shrinks, the emissions profile improves, and the operator gains a second income stream layered on top of oil production, all without waiting for a pipeline that the economics may never justify. This remains an illustrative scenario rather than a specific Flux Core case study, but every element of it is grounded in the platform documented capabilities and in the public flaring data (Source 1, Source 2).

How Flux Core Solves It

Flux Core deploys the containerized Nexus R1000 directly at the wellhead and converts stranded and associated gas into compute and token output on site. Because the platform is power agnostic, gas that is currently flared, vented, or heavily discounted becomes fuel for on site generation that powers high margin compute. No midstream takeaway is required, because the demand is created at the source rather than transported to it. The economics are direct: a molecule with negative or zero value becomes a new revenue line off an existing asset, while the flare shrinks and the emissions profile improves.

Water scarcity is the norm in producing basins, which is exactly why the closed loop liquid cooling system uses zero outside water. The fluid recirculates with no cooling tower and no evaporative loss, so there is no make up water permit and no dependence on a scarce local supply. In arid basins, this is not a convenience. It is the difference between a deployable project and an impossible one. The unit trucks to a remote pad, energizes on wellhead gas today, and retains a path to renewables or grid power later as conditions change.

Relocatability is the final piece. Fields are dynamic. Production curves decline, new pads come online, and takeaway constraints shift over time. Because the Nexus R1000 is containerized and relocatable, the deployment moves with the field. When one pad depletes, the unit relocates to the next stranded gas opportunity, protecting the operator investment and keeping the compute productive. The platform delivers 1MW per unit at Tier

III design reliability, deployable in three to six months, so an operator can turn a flare into compute revenue within a single fiscal year.

For the operator, the strategic logic is compelling on every axis at once. The financial axis adds a new revenue line from a molecule that previously had negative value. The environmental axis reduces a flare that regulators, investors, and lenders increasingly penalize, improving the emissions profile that shows up in disclosures. The operational axis avoids the cost and delay of building midstream infrastructure that may never pay back. And because the platform is power agnostic and water free, none of this depends on scarce local resources or fragile utility connections. Against a backdrop of 167 billion cubic meters flared and 54 billion dollars wasted in 2025 (Source 1), turning the flare into compute is not a marginal optimization. It is one of the largest untapped value opportunities in the entire energy sector.

The approach also scales with the operator portfolio rather than against it. A large producer typically has flaring occurring across many pads and fields simultaneously, each with its own volume, its own decline curve, and its own takeaway situation. Because Nexus units are containerized, standardized, and relocatable, a producer can deploy them progressively across the portfolio, starting with the pads where flaring is highest and takeaway is least likely, then redeploying units as fields mature. This turns a scattered, persistent liability into a managed program of on site monetization. Rather than treating flaring as an unavoidable cost of doing business, the operator can systematically convert it into compute capacity, aligning environmental performance and revenue generation across the entire asset base at once.

Conclusion and Call to Action

With 167 billion cubic meters flared and 54 billion dollars wasted in 2025 (Source 1, Source 2), stranded gas is the largest underused energy resource in the AI era. Flux Core turns it into compute at the wellhead, with a power agnostic platform, zero water closed loop cooling built for arid basins, and a relocatable form factor that follows the field. If you are flaring gas you cannot move, Flux Core turns it into compute. Let us scope a wellhead deployment on your acreage. Own the Energy, Monetize the Data.

Sources

1. Energy Connects, global gas flaring rising twice as fast as oil production (World Bank).
<https://www.energyconnects.com/opinion/features/2026/june/global-gas-flaring-rising-twice-as-fast-as-oil-production-world-bank-says/>
2. Payne Institute, Colorado School of Mines, World Bank 2026 Global Gas Flaring Tracker.
<https://payneinstitute.mines.edu/the-world-bank-2026-global-gas-flaring-tracker-report-shows-a-rise-in-flaring-for-the-third-year-in-a-row/>

5. Zero Carbon Footprint: Engineering Reductions, Not Buying Offsets

Executive Summary

A data center can advertise renewable energy credits and still burn fossil power around the clock. That gap between accounting and reality is the central problem in low carbon compute. Flux Core Data Systems and Energy rejects offset based claims and engineers real reductions into every deployment, step by step. This paper quantifies the scale of data center energy demand, explains why purchased credits do not equal real reductions, and details a four part method that combines renewable first generation, stranded and flared gas remediation, carbon credits for the true residual, and closed loop zero water cooling. The result is a low carbon posture that holds up under an actual audit rather than a marketing review.

The Problem

The energy footprint of computing is on a steep climb. The International Energy Agency projects that global data center electricity consumption will reach roughly 945 terawatt hours by 2030, about double the current level and close to 3 percent of global electricity, as reported by DataCenterDynamics (Source 1). That growth is driven overwhelmingly by AI. The question is not whether compute will use more energy. It will. The question is whether that energy is genuinely low carbon or merely labeled that way.

The common answer, purchasing renewable energy credits, does not withstand scrutiny. A credit is a financial instrument that lets a buyer claim a megawatt hour of renewable generation somewhere on the grid. It does not change what physically powers the facility. A data center can hold a portfolio of credits and still run on fossil fired grid power every hour of every day. Credits are not the same as real reductions, and increasingly, sophisticated buyers, auditors, and regulators know the difference.

The timing mismatch makes the accounting even weaker. A common practice is to match annual renewable credit purchases against annual consumption, which quietly ignores the fact that a data center runs every hour, including the many hours when the renewable generation the credit represents was not actually flowing. A facility can be fully credited on paper while drawing fossil power through every night and every windless, sunless stretch. Hourly matched, physically delivered clean power is a far higher and far rarer standard, and

it is one that a purchased credit portfolio, however large, does not meet. As the IEA projection of 945 terawatt hours by 2030 makes clear (Source 1), the absolute volume of energy involved is now large enough that this distinction has real climate consequences, not merely reporting ones.

There is also a water dimension that offset accounting ignores entirely. Data center cooling consumes enormous volumes of water. Lawrence Berkeley National Laboratory estimated roughly 17.4 billion gallons of direct water consumption by United States data centers in 2023, and the industry is under pressure to shift from evaporative cooling toward closed loop systems (Source 4). No amount of purchased carbon credit addresses the water a cooling tower evaporates. A credible environmental claim has to engineer out both the carbon and the water at the source.

The distinction between avoiding emissions and offsetting them is not academic. An offset represents a claim that a reduction happened somewhere else, often verified loosely and sometimes not additional at all, meaning the reduction might have occurred regardless. An avoided emission, by contrast, is a physical fact about the facility itself. When a deployment runs on solar rather than fossil grid power, or on gas that would otherwise have been flared, the emissions reduction is happening at the point of consumption and is directly measurable. As scrutiny of corporate climate claims sharpens, the market is moving decisively toward reductions that can be demonstrated physically rather than purchased on paper, and infrastructure that cannot make that demonstration is increasingly exposed.

Real-World Example

The direction of the industry validates the water free approach. In June 2026, NVIDIA introduced a data center design centered on closed loop liquid cooling that achieves near zero or zero on site water use, as reported by Fortune (Source 2) and analyzed by TechCrunch (Source 3). When the dominant supplier of AI hardware moves to closed loop, water free cooling, it confirms that the approach Flux Core has built around is where the industry is heading, not a fringe bet.

Consider two illustrative scenarios that show the method in practice. In the first, a site with strong solar resource runs Nexus units on solar paired with battery energy storage, displacing fossil grid power directly during the day and drawing on stored energy at night. In the second, a site near a producing oil field runs on stranded gas that would otherwise be

flared, remediating emissions that are already occurring rather than adding new load to the grid. Both scenarios are illustrative and do not describe specific Flux Core customers. Together they show how generation choice, not offset purchasing, drives the actual carbon outcome, and the flare remediation case ties directly to the World Bank finding that 167 billion cubic meters of gas were flared in 2025 (Source 5).

The two scenarios also illustrate that a low carbon posture is site specific rather than one size fits all. A location with abundant sun and land favors solar paired with storage, while a location adjacent to stranded gas favors flare remediation. A rigid design that insists on a single approach will be suboptimal in most places. Because the Flux Core platform is power agnostic, the generation strategy can be matched to the resource that delivers the largest real reduction at each site, which is precisely why engineering the outcome beats buying a uniform offset product that ignores local conditions entirely.

How Flux Core Solves It

Flux Core engineers reductions in a defined sequence rather than buying them after the fact. Step one is generation. Where the site allows, Nexus units run on renewable solar paired with battery energy storage, displacing fossil power at the source. This is a physical change to what powers the compute, not a paper claim.

Step two addresses sites where gas is needed. Flux Core prioritizes stranded and associated gas that would otherwise be flared or vented. Because that gas is already being burned with zero economic return, using it for compute reduces emissions that are already happening rather than adding new load to the grid. The World Bank data on 167 billion cubic meters of annual flaring shows how large this remediation opportunity is (Source 5).

Step three closes what remains. Only after generation choices have minimized the physical footprint does Flux Core match avoided flaring value and carbon credit mechanisms against the true residual. Credits are used to finish the job on a small remainder, not to paper over an unaddressed fossil baseline. That ordering is what separates a real reduction from an accounting claim.

Step four is cooling. The closed loop liquid system uses zero outside water, removing evaporative loss and the parasitic draw of air cooling. That holds PUE low, so more of every watt reaches the GPUs, and it takes the facility off local water supply entirely. As LBNL data

on 17.4 billion gallons of direct water use makes clear (Source 4), and as NVIDIA validated in June 2026 (Source 2, Source 3), water free cooling is now central to any credible environmental posture. Renewable first generation, flare remediation, carbon credits for the residual, and water free cooling are built in rather than bolted on afterward.

The order of these steps is the whole argument. Most low carbon claims run the sequence backward: they build a conventional fossil powered, water cooled facility and then purchase enough credits to paper over the footprint. Flux Core inverts that. It minimizes the physical footprint first through generation and cooling choices, and only then applies credits to a small, honest residual. The difference is auditable. An auditor can inspect the generation source, measure the water draw, and verify the flare reduction, none of which is possible with a facility whose green claim rests entirely on a portfolio of purchased instruments. With data center electricity demand heading toward 945 terawatt hours by 2030 (Source 1), the industry needs claims that survive that kind of inspection, and engineered reductions are the only ones that do.

Conclusion and Call to Action

With data center electricity demand heading toward 945 terawatt hours by 2030 (Source 1), the industry cannot offset its way to sustainability. Flux Core engineers the reductions instead, through renewable first generation, flared gas remediation tied to the World Bank data (Source 5), carbon credits for the true residual, and a zero water closed loop cooling system validated by the industry itself (Source 2, Source 3, Source 4). If you need infrastructure that survives a real carbon audit, ask Flux Core to walk the stack. Own the Energy, Monetize the Data.

Sources

1. DataCenterDynamics, IEA data center energy consumption set to double by 2030 to 945 TWh. <https://www.datacenterdynamics.com/en/news/iea-data-center-energy-consumption-set-to-double-by-2030-to-945twh/>
2. Fortune, NVIDIA new data center design and the AI water problem. <https://fortune.com/2026/06/22/nvidia-new-data-center-design-ai-water-problem-cooling/>
3. TechCrunch, NVIDIA wants to cut data center water use. <https://techcrunch.com/2026/06/22/nvidia-wants-to-cut-data-center-water-use-but-thats-not-the-same-as-fixing-ais-water-problem/>
4. ITIF, the data center water problem is soluble (LBNL 17.4 billion gallons, 2023). <https://itif.org/publications/2026/07/06/the-data-center-water-problem-is-soluble/>

5. Energy Connects, global gas flaring 167 bcm in 2025 (World Bank).

<https://www.energyconnects.com/opinion/features/2026/june/global-gas-flaring-rising-twice-as-fast-as-oil-production-world-bank-says/>

6. Sovereign AI, Sovereign Data: Compute That Stays Inside the Perimeter

Executive Summary

Governments and regulated institutions increasingly require that sensitive data and national AI models stay inside their own borders and under their own control. Public hyperscale cloud, by design, cannot satisfy strict residency, classification, and air gap requirements. Flux Core Data Systems and Energy places dedicated GPU compute on the customer own ground, inside a perimeter the customer owns from end to end. This paper surveys the rising tide of data localization and sovereign AI mandates, explains why shared cloud cannot meet them, and details how self contained, deployable Nexus units satisfy residency and air gap requirements by architecture, powered by a power agnostic energy path and cooled without a drop of outside water.

The Problem

Data sovereignty is moving from principle to law. Russia has revised its artificial intelligence regulation to keep sovereign models central to its framework, with the relevant law taking effect on September 1, 2026, as reported by Meduza (Source 2). The revision keeps sovereign model requirements in place even as it allows training on non Russian data, underscoring how seriously states now treat control over AI systems. Russia is not alone. Canada is pursuing a sovereign AI strategy that RBC frames as central to the country next digital chapter, emphasizing domestic compute capacity and control over national data (Source 1). Across continents, the direction is the same: keep the data, the models, and the compute inside the jurisdiction.

Public hyperscale cloud cannot answer this mandate. Its economic model is built on shared, multi tenant infrastructure spread across regions for efficiency and resilience. That architecture is fundamentally at odds with strict residency, where data must physically remain within a defined boundary; with classification, where workloads must be isolated to a controlled environment; and with air gap requirements, where systems must be physically disconnected from public networks. A contractual assurance that data will stay in a region is not the same as a physically controlled perimeter, and for defense, intelligence, and highly regulated sectors, the difference is disqualifying.

There is a further dimension that legal teams increasingly focus on: the reach of foreign law over a provider. Even when a hyperscaler operates a data center inside a given country, the provider corporate parent may be subject to the laws of another jurisdiction, which can compel disclosure of data regardless of where it physically sits. For a government handling classified information, or a bank handling regulated customer data, that exposure is unacceptable. Physical residency inside the border is necessary but not sufficient. What sovereignty ultimately requires is that the organization itself, not a foreign controlled vendor, holds exclusive physical and legal control over the hardware. That standard cannot be met by renting capacity in someone else shared facility, no matter how the contract is worded.

There is also an infrastructure problem underneath the legal one. Sovereign compute is often needed in places where the supporting infrastructure is weak: remote government sites, forward austere locations, and regions with unreliable grids and scarce water. A sovereignty solution that depends on a robust utility connection and an ample water supply cannot deploy where sovereign compute is most needed. Independence of data requires independence of power and cooling as well. A sovereign model that runs on infrastructure the sovereign does not control is sovereign in name only.

The trend toward sovereignty is also broader than any single statute. Nations that see AI as strategic infrastructure are unwilling to depend on compute controlled by foreign providers, subject to foreign law, and reachable by foreign governments. RBC frames sovereign AI as central to Canada economic and digital future, emphasizing domestic capacity and control rather than reliance on offshore hyperscale (Source 1). The same logic drives comparable initiatives across Europe, the Middle East, and Asia. For enterprises operating in regulated sectors such as banking, healthcare, and defense contracting, these national policies cascade into contractual and compliance obligations that make on soil, controlled compute a procurement requirement rather than a preference.

Real-World Example

The legal trend is concrete and dated. With Russia sovereign AI law taking effect September 1, 2026 (Source 2) and Canada actively building sovereign AI capacity (Source 1), the demand for on soil, controlled compute is no longer hypothetical. Organizations subject to these frameworks must find infrastructure that satisfies the law by design.

Consider an illustrative scenario. A defense agency or a heavily regulated national bank needs to train and run models on classified or personally sensitive data that, by law and policy, cannot leave the country or touch a shared public network. Routing that workload through a hyperscale cloud is not an option, regardless of contractual promises. This scenario is illustrative and does not describe a specific Flux Core customer, but it reflects exactly the requirements that emerging sovereign AI laws impose (Source 1, Source 2). The organization needs compute it physically controls, on ground it owns, disconnected from public networks when required.

Now consider the deployment constraint layered on top. Suppose the required site is a secure facility in a remote region with an unreliable grid and no spare water allocation, a common reality for defense and national infrastructure. A conventional data center cannot be built there on any reasonable timeline, and a hyperscale region certainly does not exist there. The organization is caught between a legal mandate to keep compute on soil and a physical reality that makes conventional on soil compute impractical. This illustrative scenario, again not a specific Flux Core case study, is exactly the intersection of legal and infrastructure constraints that a self contained, power agnostic, water free deployment is designed to resolve. The same pattern recurs for a regulated bank that must keep customer data in country, or a research ministry that must protect national datasets, wherever the required location lacks the grid and water a conventional facility assumes.

How Flux Core Solves It

Flux Core satisfies sovereignty by architecture rather than by contract. Each Nexus R1000 unit is self contained and deployable, so dedicated GPU compute can stand up inside a secure facility, a research enclave, or a forward austere site without routing sensitive workloads through a shared cloud. Residency is satisfied because the hardware sits physically on the customer ground within the required boundary. Classification and isolation are satisfied because the environment is dedicated to a single customer. Air gap requirements are satisfied because the unit can operate physically disconnected from public networks. These are properties of the deployment itself, not promises in a service agreement.

Data independence rests on power independence. Because the platform is power agnostic, it runs on stranded or associated natural gas, on solar paired with battery energy storage, on a hybrid microgrid, or on local grid power. The mission never depends on a fragile or

unavailable utility tie, which is what makes sovereign compute feasible at remote and austere sites where infrastructure is scarce. This directly addresses the deployment problem that pure grid dependent solutions cannot solve.

Cooling completes the picture. The closed loop liquid system uses zero outside water, so dense clusters can run where water and infrastructure are scarce, with no cooling tower, no make up water permit, and no dependence on a local supply. The Nexus platform delivers 1MW per unit at Tier III design reliability, targeting 99.982 percent uptime, deployable in three to six months. Sovereignty, in the Flux Core model, means owning the compute, the data, and the power as a single integrated whole.

The strategic significance of this integration is easy to underestimate. A sovereignty solution that solves data residency but depends on a foreign supplied grid connection, or on water drawn from a contested local source, has simply relocated the dependency rather than eliminated it. True sovereignty means that no external party, whether a foreign cloud provider, a local utility, or a water authority, holds a lever over the mission. By combining physically controlled compute with a power agnostic energy path and water free cooling, Flux Core removes those external levers one by one. That is why the model suits not only formal government and defense requirements but also the growing set of enterprises in banking, healthcare, and critical infrastructure whose regulators now treat compute dependence as a national security concern (Source 1, Source 2).

Conclusion and Call to Action

As sovereign AI laws take effect, from Russia framework in September 2026 (Source 2) to Canada national strategy (Source 1), organizations need compute that stays inside the perimeter by design. Flux Core delivers self contained, deployable units that satisfy residency, classification, and air gap requirements, powered by a power agnostic energy path and cooled without outside water. Talk to Flux Core about standing up sovereign AI on your terms. Own the Energy, Monetize the Data.

Sources

1. RBC, sovereign AI shaping Canada next digital chapter. <https://www.rbc.com/en/thought-leadership/the-growth-project/sovereign-ai-shaping-canadas-next-digital-chapter/>
2. Meduza, Russia revises AI regulation bill, sovereign models remain (law effective Sept 1, 2026). <https://meduza.io/en/news/2026/06/22/russia-revises-ai-regulation-bill-sovereign-models-remain-but-can-now-be-trained-on-non-russian-data>

7. College and University Data: Research Compute Without the Grid Wait

Executive Summary

Shared cloud credits run out in the middle of a training run, and a grant will not cover a re-run at on demand rates. Meanwhile, the utility upgrade a campus needs to power serious AI research is years away. That is the trap higher education keeps falling into. Flux Core Data Systems and Energy brings dedicated research compute directly to campus, up to 1MW in three to six months, without waiting on a multi year power upgrade. This paper explains why campuses cannot power their AI research ambitions on the current path, quantifies the grid wait, and details how a containerized, sovereign, power agnostic, zero water Nexus deployment lets universities own their research compute and their sensitive data at once.

The Problem

The gap between research ambition and available power is widening fast. Data center power demand is on track to double, with United States demand projected to rise toward 66 gigawatts by 2027 according to Goldman Sachs figures (Source 2), and the grid, not generation, is the binding constraint (Source 2). For a university, that macro trend translates into a very local problem: the campus substation cannot supply a serious AI cluster, and the utility upgrade to fix it is far away.

The interconnection data explains why the wait is so long. Roughly 2,600 gigawatts of capacity sits in United States interconnection queues, with a median wait of about five years to connect (Source 1). A university that wins a competitive research grant for AI work cannot tell its funders and faculty that the compute will arrive in five years. Research does not operate on that timeline, and neither do grant deliverables or graduate student careers.

Campuses also face a physical constraint that is easy to overlook. A serious AI cluster can require more power than an entire academic building, and the campus electrical infrastructure, from the substation to the distribution feeders, was never sized for that kind of concentrated load. Upgrading it is not a matter of running a new circuit. It often means a substation expansion coordinated with the local utility, which reintroduces the same multi year interconnection wait that constrains commercial developers. The research computing director who simply wants to host a cluster for a funded project finds that the request has

become a multi year capital infrastructure project, entangled with utility planning and campus master planning, long before a single GPU can be switched on.

The fallback, shared public cloud, has its own failure modes. Cloud credits are finite and expensive at scale. A long training run can exhaust an allocation partway through, and re running at on demand rates can blow through a grant budget. Just as important, public cloud raises data governance concerns. IRB governed research, sensitive human subjects data, and proprietary results are difficult to reconcile with metering hours on shared multi tenant infrastructure. Campuses need capacity they control, on ground they own, at a cost they can plan around.

The competitive stakes for institutions are rising sharply. Federal and philanthropic funding for AI research increasingly assumes access to substantial dedicated compute, and faculty recruitment and retention now hinge on whether a university can offer the resources leading researchers expect. An institution that cannot provide reliable, sovereign compute risks losing grants, talent, and standing to peers that can. The macro demand picture, with United States data center power heading toward 66 gigawatts by 2027 (Source 2), guarantees that the competition for scarce grid capacity will only intensify, which means campuses that wait for a utility upgrade are falling behind precisely when the stakes are highest.

Real-World Example

The constraints compound. With data center demand doubling toward 66 gigawatts by 2027 (Source 2) and interconnection queues imposing a five year median wait (Source 1), a campus that depends on a utility upgrade has no realistic path to timely research compute.

Consider an illustrative scenario. A research university is awarded a substantial grant for AI model development, but the campus has no power path to host the cluster the work requires. The utility upgrade is years out, and the shared cloud alternative would consume the grant in on demand fees while raising IRB and data residency concerns. This scenario is illustrative and does not describe a specific Flux Core deployment, but it is the exact bind the grid and cloud data predict for research institutions (Source 1, Source 2). The university needs dedicated compute on campus, now, under its own governance.

How Flux Core Solves It

Flux Core delivers a containerized Nexus R1000 deployment providing up to 1MW of GPU capacity for high performance computing, model training, and research computing, sited on institutional ground and live in three to six months rather than the five year grid wait (Source 1). Research computing directors keep sovereign control of sensitive datasets, IRB governed work, and grant deliverables, instead of metering hours against a public cloud bill. The data stays on campus, under campus governance, which resolves both the cost problem and the data governance problem at once.

The deployment fits how universities actually operate. Because the platform is power agnostic, the cluster can draw on existing grid capacity where it is available, on solar paired with battery energy storage, or on site generation, aligning with tight capital budgets and campus sustainability pledges alike. A university does not have to choose between its research ambitions and its climate commitments. The power path can be matched to whichever resource the campus prefers.

Closed loop liquid cooling with zero outside water keeps both utility costs and environmental reporting clean. There is no cooling tower, no make up water permit, and no draw on municipal water, which matters for institutions that publish sustainability metrics and answer to trustees and students on environmental performance. The Nexus platform runs at Tier III design reliability, targeting 99.982 percent uptime, and its containerized, relocatable form factor means the investment is not stranded if research priorities or campus master plans change. High density configurations support the accelerators modern research demands, from H200 and B200 class systems to GB300 NVL72 class density.

The procurement and governance fit is deliberate. A containerized deployment can be funded from a capital grant or a shared instrumentation award and sited on land the institution already controls, which sidesteps the multi year capital project process that a permanent building would require. Research computing staff manage the cluster under existing institutional security and data governance policies, so IRB oversight, export control obligations, and data use agreements are satisfied by keeping the workload on campus rather than negotiating them into a cloud contract. For consortia or multi campus systems, the relocatable form factor even allows capacity to be repositioned as research priorities shift across institutions, which is not possible with a fixed facility.

The cost predictability is its own argument. Public cloud bills scale with usage in ways that are notoriously hard to forecast, and a single ambitious training run can consume a semester budget without warning. A dedicated on campus cluster converts that variable, unbounded expense into a fixed, plannable one, which is far easier to reconcile with grant budgets and institutional finance cycles. Faculty can run the long experiments their research actually requires without watching a meter, and administrators can budget with confidence. In a funding environment where the demand for AI compute is only accelerating alongside the doubling of data center power demand (Source 2), that predictability is what lets an institution commit to ambitious research without fear of an unbounded bill.

There is also an educational and recruiting dividend that pure cloud access cannot provide. A physical, on campus cluster becomes a teaching instrument in its own right, giving graduate students and staff hands on experience with the operation of high performance and high density AI infrastructure, skills that are in intense demand across industry and national labs. It signals to prospective faculty and students that the institution is serious about AI research and has committed real infrastructure to it, not merely a cloud budget line that can be cut. For departments competing to attract talent in a field where compute access is often the deciding factor, owning dedicated, sovereign capacity on campus is a durable competitive advantage rather than a recurring expense to be justified each year.

Finally, the model scales gracefully as a research program grows. An institution can begin with a single Nexus unit sized to an initial grant, prove the value, and then add capacity as new awards arrive, without renegotiating a utility interconnection or launching a new construction project each time. Because the units are standardized and containerized, expansion is a matter of adding modules rather than rebuilding infrastructure. That incremental path fits the episodic, grant driven nature of academic funding far better than either a large fixed building, which risks being over or under sized, or a pure cloud arrangement, which offers no lasting asset at all. The university builds durable, sovereign research capacity one funded step at a time, on its own ground and under its own control.

Conclusion and Call to Action

With research demand rising as data center power doubles toward 66 gigawatts by 2027 (Source 2) and grid interconnection imposing a five year wait (Source 1), universities cannot afford to wait on the utility. Flux Core brings up to 1MW of sovereign, campus sited research compute in three to six months, power agnostic and cooled without outside water. If your

research ambitions are outpacing your campus power and cloud budget, ask Flux Core about dedicated university AI infrastructure. Own the Energy, Monetize the Data.

Sources

1. ChargedUp, the 2,600 gigawatt interconnection queue and five year wait (LBNL).
<https://chargeduppro.com/post/2600-gigawatt-interconnection-queue-time-to-power-2026>
2. TheStreet, AI energy appetite and the electric grid (Goldman Sachs demand growth).
<https://www.thestreet.com/markets/ais-energy-appetite-is-reshaping-the-electric-grid-data-centers>

8. Powered Land Lease Rates: The Real Bubble in AI Infrastructure

Executive Summary

The loudest debate in AI infrastructure asks whether the whole buildout is a bubble. That framing misses the real risk. The bubble is not AI compute. It is powered land. Consider a recent quote: 4MW of powered land at 3 million dollars a year, with electricity billed separately on top. That works out to 750,000 dollars per megawatt per year, just to occupy land and hold a spot in an interconnection queue, before a single GPU is installed or a single token is served. This paper does the math on powered land lease rates, argues that the exposure is a commercial real estate bubble rather than an AI bubble, and explains why the rise of behind the meter power causes those rent rolls to collapse. Flux Core Data Systems and Energy stands on the other side of that trade.

The Problem

Powered land, land with a viable path to grid power, has become one of the most sought after assets in the AI economy, and landlords have priced it accordingly. The math on a recent quote is stark. A quote of 4MW of powered land at 3 million dollars per year equals 750,000 dollars per megawatt every year, with electricity billed separately on top. That is 750,000 dollars per MW per year to occupy land and hold a spot in an interconnection queue, before a single GPU is installed or a single token is served. The tenant is paying a premium not for compute, not even for power, but for a queue position and a patch of ground.

That premium exists only because of a temporary scarcity: the grid cannot connect load fast enough. Roughly 2,600 gigawatts of capacity sits in United States interconnection queues, with a median wait of about five years (Source 2). As long as a grid connection is the only way to power compute, a piece of land with a queue position commands a scarcity rent. Landlords are writing decade long leases against that scarcity, betting it will persist.

The bet is unsound, because the scarcity is already eroding. SemiAnalysis projects more than 40 gigawatts of behind the meter data center capacity in the United States by 2028 (Source 1), and S&P Global expects behind the meter generation to meet roughly 25 percent of new data center demand by 2030 (Source 3). Once a developer can generate power on site, the grid queue position that justified the 750,000 dollar per MW rent loses its

value. This is the classic setup for a commercial real estate bubble: long dated leases priced against a scarcity that a competing technology is about to eliminate.

It is worth being precise about why this is a real estate bubble and not an AI bubble, because the two are constantly conflated. Demand for AI compute is real, large, and growing, as the doubling of data center power demand demonstrates. That demand is not speculative. What is speculative is the assumption that the only way to serve it is through grid connected land, and therefore that grid connected land commands a permanent premium. Bubbles form when an asset is priced on an assumption that a structural change is about to invalidate. Powered land is priced on the assumption that the grid remains the sole gateway to power. Behind the meter generation removes that assumption, and when the assumption goes, the premium goes with it, even as the underlying AI demand keeps climbing.

Real-World Example

The example is the quote itself. At 750,000 dollars per MW per year for 4MW, a tenant pays 3 million dollars annually for land and a queue spot, with electricity on top. Set that against the behind the meter trajectory: more than 40 gigawatts by 2028 (Source 1) and about a quarter of new demand by 2030 (Source 3). When behind the meter generation becomes the default, a 750,000 dollar per MW ground lease is worth a fraction of the paper it is printed on.

Consider an illustrative scenario that shows the alternative. A developer faced with that powered land quote runs the numbers and declines. Instead of paying 750,000 dollars per MW to rent a queue position with a five year wait (Source 2), the developer chooses to generate power on site behind the meter and energize in months. This scenario is illustrative and does not describe a specific Flux Core customer, but it is precisely the choice the market data says developers are increasingly making (Source 1, Source 3). Every developer who makes that choice removes a future tenant from the powered land market, which is exactly how the rent roll collapses.

How Flux Core Solves It

Flux Core stands on the other side of the powered land trade. We understand power, so we develop our own opportunities rather than rent someone else bottleneck. There is no ground lease to pay, no interconnection queue to wait in, and no electricity pass through

markup layered on top of the rent. The 750,000 dollar per MW per year structure simply does not exist in the Flux Core model, because Flux Core is not renting scarcity. It is generating power.

The deployments are power agnostic across stranded or associated gas, renewable solar paired with battery energy storage, hybrid microgrids, and grid power where it is available. That flexibility is what makes the behind the meter strategy practical, and it is what aligns Flux Core with the market shift that SemiAnalysis and S&P Global describe rather than against it (Source 1, Source 3). Instead of betting on a scarcity that is eroding, Flux Core is helping erode it.

The Nexus R1000 delivers this as a 1MW containerized data center, deployable in three to six months, built to Tier III design standards, and cooled by a closed loop liquid system that uses zero outside water. That last point matters for siting, because a deployment free of water permits and cooling towers is not tethered to the scarce, expensive, queue burdened parcels that command powered land rents. Flux Core can site where the power is, not where the queue position happens to be, which is the structural advantage that makes the powered land bubble irrelevant to its business.

Consider what the same 4MW looks like under the Flux Core model versus the powered land quote. Under the lease, the tenant pays 3 million dollars per year for land and a queue spot, waits years for interconnection, and then pays for electricity on top with a pass through markup. Under the Flux Core model, four Nexus units are deployed on a site chosen for its power resource, energized in months, and cooled without a drop of outside water, with the energy owned and operated rather than rented at a marked up rate. The contrast is not incremental. One path pays a premium to wait in a queue. The other path skips the queue entirely by making power where the compute sits, which is the whole point of behind the meter and the reason the land premium cannot hold.

The timing of this shift is what makes the powered land position so precarious. Decade long leases assume the scarcity persists for a decade. But the behind the meter transition is not a distant possibility. It is already underway, with SemiAnalysis forecasting more than 40 gigawatts by 2028 and S&P Global projecting a quarter of new demand by 2030 (Source 1, Source 3). A landlord signing a ten year powered land lease today is betting against a trend that credible analysts expect to reshape the market within the first few years of that lease.

When the tenants who would have paid the premium instead build behind the meter, the leases do not gracefully reprice. They are stranded, because the whole value proposition, exclusive access to grid power, has evaporated. That is the mechanism by which a commercial real estate bubble pops, and it is why Flux Core has built its entire model on generating power rather than renting access to it.

Conclusion and Call to Action

The AI buildout is not the bubble. Powered land as a service is the bubble, and it pops as behind the meter power scales past 40 gigawatts by 2028 (Source 1) and toward a quarter of new demand by 2030 (Source 3). A 750,000 dollar per MW per year ground lease is a bet on a scarcity that a five year queue (Source 2) can no longer protect. Flux Core owns the energy instead of renting the bottleneck, with a power agnostic, zero water platform and no ground lease, no queue, and no electricity markup. Stop renting a power problem. Talk to Flux Core about owning the energy and monetizing the data. Own the Energy, Monetize the Data.

Sources

1. SemiAnalysis, United States grid constraints toward 40 gigawatts behind the meter by 2028. <https://newsletter.semianalysis.com/p/us-grid-constraints-towards-40gw>
2. ChargedUp, the 2,600 gigawatt interconnection queue and five year wait. <https://chargeduppro.com/post/2600-gigawatt-interconnection-queue-time-to-power-2026>
3. S&P Global Market Intelligence, behind the meter to meet about 25 percent of new demand by 2030. <https://www.spglobal.com/market-intelligence/en/news-insights/research/2026/06/ai-data-center-power-demand-grid-constraints-energy-resilience>