

The Pluralistic Future of AI: A Comprehensive Analysis of Decentralized Large Language Models

Nursan Omarov* and Alimzahn Tokushev

Independent Researcher, Kazakhstan

*Correspondence: Nursan Omarov, Independent Researcher, Kazakhstan, E-mail: nursan2007@gmail.com;

DOI: <https://doi.org/10.56147/aaiet.1.6.81>

Citation: Omarov N, Tokushev A (2025) The Pluralistic Future of AI: A Comprehensive Analysis of Decentralized Large Language Models. J Adv Arti Int Eng & Techn 1: 81.

Abstract

The centralized status quo: The current landscape of Artificial Intelligence (AI), particularly concerning Large Language Models (LLMs), is dominated by a centralized paradigm. Industry leaders such as OpenAI, Anthropic and Google DeepMind deploy and manage their models on massive computational infrastructure within expansive data centers. This architecture, where a single, colossal model serves millions of users *via* cloud-based APIs, has enabled unprecedented scalability, high performance and the capacity for continuous updates. However, this model is not without significant drawbacks. It requires substantial financial investment in infrastructure, is entirely reliant on constant internet connectivity and introduces considerable privacy and data control concerns, as sensitive information must be processed in the cloud. The user also frequently pays for compute capacity they do not fully utilize, leading to an inefficient cost model based on subscriptions or per-token usage.

Keywords: AI: Artificial Intelligence; LLM: Large Language Model; Centralized LLM; Cloud-based large language model; Decentralized LLM; Locally deployed large language model; Data centers; Model; API

Received date: September 09, 2025; **Accepted date:** September 29, 2025; **Published date:** September 22, 2025

Introduction

The centralized status quo

The current landscape of Artificial Intelligence (AI), particularly concerning Large Language Models (LLMs), is dominated by a centralized paradigm. Industry leaders such as OpenAI, Anthropic and Google DeepMind deploy and manage their models on massive computational infrastructure within expansive data centers [1]. This architecture, where a single, colossal model serves millions of users *via* cloud-based APIs, has enabled unprecedented scalability, high performance and the capacity for continuous updates [1].

However, this model is not without significant drawbacks. It requires substantial financial investment in infrastructure, is entirely reliant on constant internet connectivity and introduces considerable privacy and data control concerns, as sensitive information must be processed in the cloud [1]. The user also frequently pays for compute capacity they do not fully utilize, leading to an

inefficient cost model based on subscriptions or per-token usage [1].

The decentralized alternative and the pluralistic thesis

An alternative approach is emerging: decentralized LLMs, where smaller, optimized models are deployed and executed locally on end-user devices, such as PCs, mobiles and edge hardware [1]. This paradigm represents a fundamental inversion of the centralized model, reminiscent of the 20th century transition from centralized mainframe computing to personal computers [1]. The core argument of this report is that the future of AI is not a zero-sum competition between these two models but a "pluralistic" one, where both centralized and decentralized architectures coexist and serve distinct, complementary needs [1]. Centralized systems may continue to dominate for large-scale, complex reasoning, while decentralized systems will enable a new era of personalized, private and resilient intelligence for individuals and organizations [1].

Scope and structure of the report

This report will first provide a foundational comparative analysis of the centralized and decentralized LLM paradigms. Following this, it will delve into the technical enablers that make on-device AI a reality, exploring model compression, optimization and distributed training frameworks. The analysis will then extend to the economic, environmental and ethical implications of this shift. Finally, the report will present a series of domain-specific applications and case studies to illustrate the real-world impact of decentralized AI and conclude with an outlook on the future research trajectory.

Foundational Concepts: Centralized vs. Decentralized LLM Architectures

The centralized paradigm and its bottlenecks

The architecture of centralized LLMs is defined by its scale and unified deployment. With parameter counts reaching into the trillions, these models are housed in cloud servers and are capable of complex reasoning across diverse domains [1]. Their advantages are clear: They offer high performance and generality, benefit from continuous updates and improvements and can handle a wide array of tasks [1].

However, this model creates a performance-to-accessibility bottleneck. The immense infrastructure investment required restricts their deployment to a handful of major corporations, creating a centralized control that can stifle innovation and limit user agency [1]. Furthermore, their dependence on continuous internet connectivity makes them unsuitable for applications in remote areas, disaster recovery scenarios or any environment with limited bandwidth [1]. The core trade-off here is that centralized systems offer unparalleled scale and general capabilities but are constrained by their reliance on vast, expensive and internet-dependent infrastructure.

The decentralized alternative: On-device LLMs

In contrast, the decentralized paradigm deploys smaller, optimized models directly onto end-user devices [1,2]. The fundamental benefits of this architecture are multifold. It offers significant cost-efficiency, as users incur a one-time device-level cost rather than recurring subscription or per-token fees [1]. This approach also provides robust privacy, as sensitive data remains on the device and is not processed in an external cloud, mitigating the risk of data leaks or hacks [3]. The offline functionality of these models is critical for resilience and autonomy in environments with limited or no connectivity [1]. Finally, on-device models can be fine-tuned to narrow, user-specific tasks, enabling a high degree of personalization [1]. The fundamental trade-off of this paradigm is a sacrifice of raw, general-purpose performance for

ubiquity, privacy and cost-efficiency.

Comparative analysis of paradigms

The core dynamic driving the evolution of LLMs is the nuanced trade-off between the scale of centralized models and the accessibility of decentralized ones. The following table expands on the initial user document to provide a comprehensive comparison across several key dimensions, illustrating how these two paradigms cater to different use cases and constraints (**Table 1**).

Table 1: Trade-off between the scale of centralized models and the accessibility of decentralized ones.

Feature	Centralized LLMs	Decentralized LLMs
Scale	Trillions of parameters	Millions-billions (optimized)
Deployment	Cloud servers in data centres	Local devices (PCs, mobiles, edge)
Connectivity	Requires continuous internet	Fully offline possible
Cost Model	Subscription/per-token usage	One-time, device-level cost
Privacy	Data processed in the cloud	Data remains local
Personalization	Limited (general-purpose)	High (task-specific tuning)
Governance	Centralized (corporate control)	Distributed (user/community control)

The Technical Enablers of On-Device AI

The ability to run large language models on resource-constrained devices is not a simple feat but the result of significant advancements in model compression and distributed inference. These technical innovations are the foundation upon which the decentralized AI revolution is built.

Model compression and optimization

The primary challenge for on-device deployment is the sheer size and computational demand of LLMs. Model compression techniques address this by reducing memory footprint and computational requirements while minimizing performance degradation.

Quantization

Quantization is a key technique that reduces the numerical precision of a model's weights and activations, typically from 32-bit floating-point to lower bit-width representations like 8-bit integers [5]. A straightforward approach is Post-Training Quantization (PTQ), which applies quantization after the model has been trained. This is simple to implement but can lead to performance drops for complex models [5]. A more advanced method is Quantization-Aware Training (QAT), which incorporates the effects of quantization during the training process

itself, allowing the model to adapt to reduced precision and maintain higher accuracy. The evolution from simple PTQ to more sophisticated techniques such as outlier-aware quantization and mixed-precision quantization demonstrates that the technical challenge is not merely reducing size but doing so without critically compromising performance [5]. This is an active area of research aimed at bridging the gap between massive models and limited hardware.

Pruning

Pruning is the process of eliminating redundant or unnecessary parameters from a neural network, thereby reducing its size and computational requirements. Research has shown that up to half of a model's parameters can be removed with almost no impact on performance, highlighting the common issue of over-parameterization in large models [6]. By identifying and trimming these excess components, pruning makes models more efficient for on-device deployment.

Knowledge distillation

Knowledge distillation is a powerful technique where a large, high-performing teacher model transfers its knowledge to a smaller, more efficient "student" model. The student model is trained to mimic the outputs of the teacher, effectively learning the same generalized knowledge but with a significantly smaller parameter count. A prominent example is DistilBERT, which achieves a 50% reduction in model size while retaining near-equivalent performance on most tasks [6]. This process represents a direct, symbiotic link between the centralized and decentralized paradigms. The expensive, massive-scale training of a centralized model (the teacher) is leveraged to produce a lightweight, efficient on-device model (the student), perfectly illustrating how one paradigm can feed and enhance the other in a pluralistic ecosystem.

Distributed inference architectures

For models that are still too large to fit on a single device, even after compression, distributed inference offers a solution. Tensor parallelism partitions the model's neural network tensors, such as weight matrices, across multiple edge devices for collaborative inference [7]. While this enables the deployment of larger models, it introduces the significant challenge of communication overhead due to frequent all-reduce operations needed to aggregate intermediate outputs across devices [7].

The Distributed Paradigm: Training and Synchronization

The principles of decentralization extend beyond model deployment to the very process of training and model synchronization. This distributed approach addresses core challenges related to privacy, resource

constraints and development bottlenecks.

Federated Learning (FL) for privacy-preserving training

Federated Learning (FL) is a decentralized AI strategy that allows a global model to be trained across a multitude of devices without requiring the raw data to be centralized or shared [8]. The framework operates with a central server or aggregator that sends out the latest model to participating devices or workers [8]. Each worker then uses its own local data to train and update the model. Crucially, only the model updates not the sensitive data are sent back to the aggregator which then combines these updates to refine the global model [8]. This approach harnesses the collective computational power of distributed devices, reducing the burden on a central server and minimizing data transmission costs [8].

While FL is a significant step toward privacy-preserving training, it is not a complete solution. The framework introduces new security challenges, as adversaries may exploit the shared gradients during training to infer sensitive information about the underlying data. This necessitates the use of advanced cryptographic techniques like Secure Multi-Party Computation (SMPC) and homomorphic encryption to safeguard against such threats [8]. This demonstrates that FL, while solving one problem, introduces its own set of complex research challenges that must be addressed for it to be a viable long-term solution [9].

Advanced distributed architectures

Decentralized Mixture of Experts (MoE)

A Mixture of Experts (MoE) is an architecture where a router network directs input tokens to a sparse subset of specialized "expert" sub-networks [11,12]. The centralized paradigm is often constrained by the need for a massive, high-bandwidth network fabric to synchronize gradients across thousands of GPUs during training [13]. MoE offers an elegant solution by providing an orthogonal form of parallelism that can be applied in a decentralized manner [13]. Instead of synchronizing gradients, the training burden is partitioned across independent expert models, each trained on its own compute island with no cross-communication [13]. This distributed approach to development allows for the use of scattered, heterogeneous hardware, democratizing the creation of massive models and alleviating the systems constraints that limit centralized training runs [13].

Federated MoE frameworks

The integration of MoE with federated learning creates a sophisticated, hybrid system. Frameworks like Federated Mixture of Experts (FedMix) and FedMoE-DA allow for the training of an ensemble of specialized models within an FL setup [14,15]. This architecture leverages the diversity of local client data to train specialized experts, enhancing

both robustness and persuasibility while maintaining privacy [15]. The combination of FL, which provides a mechanism for privacy-preserving generalization and MoE, which provides a mechanism for privacy-preserving specialization, creates a powerful system that can collectively learn from diverse data sources while maintaining individual data control and fostering domain-specific expertise [16].

Economic and Environmental Implications

Economic democratization of AI

The shift toward decentralized LLMs has profound economic implications. By eliminating recurring API costs, it significantly lowers the barrier to entry, allowing individuals and small businesses to leverage advanced AI without relying on corporate APIs [1]. This shift from a B2C (Business-to-Consumer) API-as-a-service model to a P2P (Peer-to-Peer) marketplace model is creating new ecosystems [2]. Platforms like SingularityNET and Bittensor are emerging as decentralized marketplaces where AI models and datasets can be bought, sold and collaboratively developed [2]. Autonomous AI agents are already transforming the Web3 economy by optimizing trades, managing liquidity pools and executing financial operations without human oversight [2].

However, the decentralized AI economy is still in its nascent stages and is subject to considerable hype and misinformation [17,18]. The presence of fraudulent projects and scams that lead to rug pulls is a significant concern [2]. This highlights the need for a balanced perspective that acknowledges the immense potential of this new economic model while also recognizing its speculative and risky nature.

The AI energy challenge

The widespread adoption of AI is not sustainable under the centralized paradigm due to its immense energy footprint. Data centers, primarily powered by energy-intensive GPUs, already consume more electricity than entire nations and are projected to double their energy use to 500 TWh by 2027 [19]. A single query on a centralized model can consume approximately 2.9 Wh of electricity, roughly 10 times more than a standard Google search [19]. The cumulative energy consumption from continuous inference across millions of users far exceeds the energy used in a single training run [20].

On-device AI presents a crucial solution to this environmental crisis. By processing data locally, it eliminates the need for energy-intensive data transmission to and from distant data centers [19]. The use of specialized, energy-efficient chips for local processing results in a dramatic 100-1000-fold reduction in energy consumption per task compared to cloud-based AI [19]. The technical enablers of model compression discussed

earlier are directly linked to this environmental benefit, making the push for on-device AI not just a technological refinement but a strategic imperative for the industry's sustainability.

Ethical and Societal Considerations

Privacy and data confidentiality

The most prominent ethical benefit of decentralized LLMs is their ability to preserve privacy. When a model runs locally, sensitive professional or personal data never leaves the device, ensuring confidentiality [3]. This stands in stark contrast to centralized models, where corporate policies regarding data retention and privacy can change overnight and often do not guarantee the confidentiality of data submitted through chatbot interfaces, even for paid subscription plans [4].

The challenge of bias and validation

While local LLMs offer a privacy advantage, they are not immune to ethical challenges. They inherit biases including gender, racial and socio-economic prejudices from the foundational datasets on which they were trained [21]. These biases can be inadvertently perpetuated and amplified in the model's outputs [21].

Furthermore, the power of local fine-tuning, while a tool for personalization, also carries the risk of introducing or reinforcing new biases based on the user's specific data [23]. This raises a critical question about accountability: With a centralized model, the responsibility for a biased output can be traced to a single corporate entity. With a decentralized, locally fine-tuned model, accountability becomes distributed and ambiguous. The report suggests that a critical and unresolved ethical and legal challenge is determining who is responsible for a harmful output the foundational model developer, the fine-tuning platform or the end-user. This underscores the need for trust and validation in locally fine-tuned models, especially for sensitive applications like healthcare or law [1].

Societal polarization and echo chambers

Centralized social media platforms and their AI-driven recommendation algorithms, motivated by a profit imperative to maximize user engagement, are known to create filter bubbles and echo chambers [24]. These systems reinforce users' existing beliefs and can be instrumental in spreading misinformation and even radicalization [25].

The rise of personalized, on-device AI introduces a new dimension to this problem. A personal AI could theoretically act as an agent of change, proactively exposing a user to diverse viewpoints to counter the filter bubble effect [27]. However, it could also be fine-tuned by the user to align perfectly with their existing biases, creating a more potent, self-directed personal echo chamber that is far more difficult to govern or mitigate

[24]. The problem shifts from being an external, corporate-controlled issue to a personal, user-controlled one, which is a far more complex and nuanced ethical dilemma.

The impact on human cognition

A broader societal risk of pervasive AI is the potential for cognitive offloading, where humans outsource complex problem-solving and critical thinking to AI systems [28]. Research has already revealed a significant negative correlation between frequent AI tool usage and critical thinking abilities, suggesting that an over-reliance on these systems may come at a cognitive cost [29].

On-device AI, being always available and highly personalized, could exacerbate this trend, raising a fundamental question about the trade-off between efficiency and intellectual autonomy in a future where every person carries their own intelligent system.

Domain-Specific Applications and Case Studies

The following **table 2** provides a summary of real-world applications and case studies that highlight the practical benefits and challenges of decentralized AI across various domains.

Table 2: Summary of real-world applications and case studies.

Domain	Specific application	Benefit	Key technology	Example/case study
Healthcare	Offline diagnostic assistants for rural areas	Resilience, privacy, cost-efficiency	On-device LLMs Federated Learning (FL)	OfflineMedics, MONAI, FedMRG framework [1,33]
LegalTech	Local document validation and drafting	Privacy, cost-efficiency, speed	On-device LLMs	LexiHK [37]
Education	Personalized offline tutors and assistants	Resilience, personalization	On-device LLMs, Retrieval-Augmented Generation (RAG)	Khanmigo, Edukapi [42]
Aerospace	Offline AI for problem-solving	Resilience, autonomy	On-device LLMs	Astronauts relying on offline AI [1]

Healthcare

Decentralized AI is poised to revolutionize healthcare. On-device LLMs can act as diagnostic assistants in rural or remote areas with limited or non-existent internet connectivity, providing crucial support in emergency situations [1]. AI models can interpret medical imaging, detect bone fractures and triage patients with greater speed and accuracy than humans in many cases [30,31].

A particularly compelling application is the use of Federated Learning in medical research. This framework enables collaborative model training across multiple hospitals without the need to share sensitive patient data, which is often restricted due to privacy concerns [30]. Projects like MONAI (Medical Open Network for AI) and the FedMRG framework are demonstrating how this technology can build more robust models from diverse datasets, particularly for rare diseases, while maintaining data confidentiality [33].

LegalTech

For small law firms, where time and resources are limited, local LLMs offer significant advantages. These models can automate repetitive tasks such as document summarization, classification and initial drafting of legal documents, freeing up lawyers to focus on high-value, critical thinking tasks [34-36]. Local deployment ensures that sensitive client information never leaves the firm's devices.

A key case study demonstrating this trend is the development of LexiHK, a fine-tuned local LLM for legal

document assistance developed by the Hong Kong Department of Justice [37]. This demonstrates a policy-driven move toward the adoption of local, domain-specific models to enhance efficiency and security within the legal sector.

Education

On-device LLMs can serve as personalized, offline tutors that adapt to each student's unique learning style and pace [1]. These tools can provide real-time feedback, assist with homework and streamline administrative tasks for teachers, such as creating lesson plans and rubrics [39].

While the theoretical promise of a fully offline, private AI tutor is strong, an examination of real-world applications reveals a more nuanced reality [40,41]. For example, while Khanmigo offers engaging, on-topic tutoring, other platforms like Flexi require an internet connection and applications like Edukapi may still collect user data despite the promise of on-device functionality [42,43]. This illustrates a crucial point: The clear theoretical distinction between centralized and decentralized models is often blurred in practice, with many commercial products operating as hybrids that leverage both local and cloud-based components to deliver their services.

Conclusion and Future Directions

The report concludes that the future of AI is not a simple choice between centralized or decentralized systems but a pluralistic coexistence where each paradigm addresses unique needs and constraints. Centralized

models will continue to be the engine for large-scale, general-purpose intelligence, while decentralized models will democratize access, ensure privacy and enable resilience in a new era of on-device, personalized intelligence.

However, several challenges remain. The performance gap between on-device and centralized models, while shrinking due to advancements in compression and optimization, has not yet been fully closed. The technical complexity of distributed training, particularly in federated learning and decentralized MoE architectures, introduces new security and synchronization challenges that require ongoing research. On a societal level, the shift in ethical responsibility from a centralized corporate entity to a distributed network of users creates new legal and ethical dilemmas that are yet to be resolved. Finally, the broader societal impact on human cognition and the potential for a new form of personalized, self-directed echo chamber must be carefully navigated.

Despite these challenges, the trajectory is clear. The report reinforces the historical analogy to the personal computer revolution, where centralized mainframes gave way to a distributed computing model. A similar shift is emerging in AI, promising a future where every person can carry their own intelligent system, privately and cost-effectively, without dependence on external servers. This vision offers a more sustainable, equitable and autonomous technological future.

References

1. How will decentralized AI affect big tech?
2. Offline AI made easy: How to run large language models locally.
3. Local LLMs: Ethical, secure and sustainable AI.
4. LLM optimization: Quantization, pruning and distillation medium.
5. Quantization, distillation and pruning.
6. Distributed on-device LLM inference with over-the-air computation.
7. Toward federated large language models.
8. Federated learning: 5 use cases & real-life examples.
9. What is decentralized AI? A beginner's guide to blockchain-powered intelligence.
10. Mixture-of-Experts (MoE) models in AI.
11. Mixture of experts.
12. Decentralized diffusion models.
13. Parallelism and distributed training for maximizing AI efficiency.
14. Federated mixture of experts.
15. Federated learning using a mixture of experts.
16. AI + Blockchain: The most promising projects to watch in 2025.
17. Understanding Decentralized Finance (DeFi): Basics and functionality.
18. How on-device AI could help us to cut AI's energy demand.
19. Energy efficiency in AI models: Strategies for a sustainable future.
20. Ethical implications and challenges of using language models.
21. Large language model.
22. Detecting bias in large language models: Fine-tuned KcBERT.
23. Echo chambers and recommendation algorithms: Who decides what we see online?
24. Algorithmic radicalization.
25. We'd like to use additional cookies to understand how you use the site and improve our services. UK Parliament Committees.
26. Trapped in a social media echo chamber? A new study reveals how ai can offer an escape.
27. AI and society: Implications for global equality and quality of life.
28. AI tools in society: Impacts on cognitive offloading and the future of critical thinking.
29. Decentralized AI: What are the advantages for the healthcare industry?
30. 7 ways AI is transforming healthcare: The world economic forum.
31. Aidoc (2025) Clinical AI solutions for healthcare providers.
32. LLM-driven medical report generation *via* communication-efficient heterogeneous federated learning.
33. Federated learning for medical applications: A taxonomy, current trends, challenges and future research directions.
34. Medical open network for AI.
35. Small law firm AI guide: Using LLMs in 2025.
36. LCQ5: Application of legal technology and artificial intelligence.
37. Personalized learning with AI: Transforming education.
38. Top 10 AI-powered learning experience platforms in 2025.
39. LLM in education: The secret to smarter and personalized learning.
40. SchoolAI (2025) Reimagining student success.
41. Khan academy's AI-powered teaching assistant & tutor.
42. Edukapi: Your AI tutor 24/7 apps on Google Play.
43. Flexi: A free science and math AI tutor for every student.