

The Next Frontier of AI: Emerging Techniques and Future Trends

Venkata Siva Prasad Maddala

CVS Health, USA

Abstract

The landscape of artificial intelligence has evolved from narrow, task-specific systems into sophisticated platforms that are capable of autonomous reasoning and multi-dimensional data processing. Contemporary architectures demonstrate the ability for transformation via chain-of-thought prompting mechanisms that allow complex reasoning and foundation models exhibiting remarkable transfer learning across domains. Autonomous agentic systems constitute a paradigm shift towards proactive agents pursuing complex objectives through multi-step planning, while multimodal intelligence integrates diversity in sensory channels for comprehensive contextual understanding. The architectural shift to edge computing is placing latency, privacy, and energy efficiency on special hardware and neuromorphic computing models. Federated learning is a privacy-preserving method that enables the collaborative training of models without accumulating sensitive data, and explainable AI satisfies the needs of interpretability in situations with high-stakes decisions. By incorporating intelligent capabilities directly into the existing workflows via collaborative intelligence schemes, and the overall governance scheme to deal with accountability, fairness, and regulatory compliance. These overlapping developments will provide the basics of AI systems that are auxiliary to human capacity but congruent to moral norms and social values in healthcare, autonomous vehicles, money, and smart city applications.

Keywords: Autonomous Agents, Multimodal Intelligence, Edge Computing, Federated Learning, Algorithmic Governance

I. Introduction: The Generation of Artificial Intelligence Systems

Artificial intelligence has shifted drastically from the terrain of small and focused systems to complex systems that can apply independent reasoning and multi-dimensional data analysis. Contemporary AI machines can achieve much more than pattern recognition and prediction, but exhibit sophisticated forms of thinking that provide the ability to plan strategically, comprehend the surrounding context, and make decisions in different settings. This development is basically a paradigm shift in the architectural concepts of the intelligent systems, where conventional methodologies were fully established with a clear set of constraints, with very little adaptability or autonomy in solving problems.

Recent AI systems are also structured with sophisticated reasoning schemes that allow the systems to decompose complex goals into action subgoals, take into account multiple ways of solutions, and adapt strategies based on environmental responses. Research on chain-of-thought prompting has shown that it is possible for language models to execute complex reasoning by explicitly generating the intermediate steps of reasoning and, therefore, demonstrate emergent abilities in arithmetic reasoning, commonsense reasoning, and symbolic reasoning tasks. In fact, studies have found that asking the models to explicitly state their step-by-step reasoning process before giving a final answer significantly enhances performance in tasks involving multi-step logical inference, mathematical problem-solving, and strategic planning [1]. These are emerging properties from the scaling of model parameters and training data, indicating that with a large enough scale, models internally develop representations supportive of complex reasoning without explicit programming of logical operations.

10.48047/jocaaa.2026.35.01.68

Combinable with this intersection of various technological innovations, the key catalyst to this change has been foundation models: big models trained on general data and specialized for a large variety of downstream tasks. So, foundation models are developed to symbolize a paradigm shift, where one model, trained using a variety of data, acquires general-purpose abilities that can be applied across fields. Modelling these models shows impressive transfer learning performance, in which what they learn in the pre-training phase can be efficiently performed on tasks that had not been directly observed during training. At the same time, homogenization around the foundation model paradigm has enabled centralized governance and standardized evaluation but has also given rise to a number of concerns involving power concentration, environmental impact due to computational requirements, and systemic risks when a variety of applications rely on shared model infrastructure. This impacts virtually all sectors, from diagnostics and scientific research in healthcare to industrial automation and creative production, placing governance frameworks that balance their potential for transformation with alignment in service of human values and societal interest at the center.

II. Autonomous Agentic AI: From Task Execution to Strategic Reasoning

At that, agentic AI can be seen as a paradigm shift in respect of reactive systems and active agents that can perform multistage objectives by relying on multi-step reasoning and planning. The systems are indeed traditional in terms of agency, in which each agent views the environment, creates goals, develops strategies, implements sets of actions, and reinvents strategies as the outcomes and circumstances evolve. Most systems' architectural core combines reasoning engines for goal decomposition, a memory system that also maintains working and long-term knowledge, tool-use to interact with other external resources, and feedback loops to allow for continuous learning from outcomes.

Recent work has demonstrated that large language models can form a good basis for autonomous agents, possessing natural language understanding, knowledge memorization, and reasoning capabilities. Several recent surveys suggest that there are four key building blocks of an agent architecture: a profiling module specifying the agent's role and behaviors, a memory module storing interactions and prior knowledge, a planning module breaking complex tasks into sub-feasible subtasks, and an action module posting decisions via tool or environment interaction. The profiling component determines the agent's identity and area of expertise, enabling role-specific behaviours, such as customer service representatives, research assistants, and supply chain coordinators. Designs for the memory architecture vary on the basis of shortterm memory for context, and long-term memory storing experience that will also inform decisions at some later date, and retrieval mechanisms to determine which history is relevant to goals today [3]. Planning capabilities allow agents to build multistep plans toward high-level objectives by applying methods that range from simple decomposition schemes to complex planning using feedback loops. Agents generate plans without feedback, plans with feedback about environmental outcomes at each step, or construct complete plans that are iteratively refined based on execution results. The action module executes the plans, translating them into concrete operations such as tool invocation to access external databases or other computational resources, embodied actions in either physical or simulated environments, and communication with other agents or human supervisors. Work has shown that agents equipped with appropriate tools can successfully solve complex benchmarks that require coordination among dozens of interdependent actions over long time horizons in the presence of uncertainty and incomplete information [3].

Multi-agent systems further extend this capability through collaboration, with diverse agents coordinating their behavior to solve problems that exceed the capability of a single agent. Researchers investigating social influence as intrinsic motivation have found that cooperative behaviors can emerge from reinforcement learning mechanisms that reward agents for their influence over the behavior of other agents.

10.48047/jocaaa.2026.35.01.68

Experimental work has demonstrated that agents receive rewards for influencing teammates and learn effective communication, coordinated strategy, and emergent collaborative behaviors in the absence of explicit coordination protocols. One of the most successful applications of this approach is in partially observable environments, where agents must share information to develop comprehensive situational awareness. The intrinsic motivation framework encourages agents to take actions that maximize their causal influence on the behavior of teammates, naturally giving rise to teaching behaviors in which experienced agents guide less-experienced partners [4]. However, making systems operate autonomously presents significant challenges for governance in terms of oversight, control, and alignment, demanding robust monitoring frameworks and mechanisms for intervention that provide insight into agent reasoning and prevent harmful outcomes.

Architecture Module	Primary Function	Key Capability
Profiling Module	Role Definition	Identity establishment, behavior specification
Memory Module	Information Storage	Short-term context, long-term experience retention
Planning Module	Task Decomposition	Multi-step strategy construction, feedback integration
Action Module	Execution	Tool invocation, environmental interaction, communication
Multi-Agent Coordination	Collaborative Problem-Solving	Social influence, emergent behaviors, information sharing

Table 1: Autonomous Agent Architecture Components and Functions [3, 4]

III. Multimodal Intelligence: Coherence of Multiple Data Streams in Deep Understanding

In the case of multimodal AI systems, information is combined in the various types of sensory channels and data, which represent an enormous step forward in general machine intelligence. These different sensory modalities supply textual, visual, and auditory input that is synthesized by the new architectures to generate single representations that reflect much more contextual information and therefore induce subtle comprehension of complex situations. Their technical backbone depends on sophisticated encoding mechanisms that transform various data types into compatible representational spaces in which information from different modalities can be meaningfully combined through cross-modal attention mechanisms and fusion strategies.

Comprehensive analysis of multimodal machine learning identifies six core technical challenges spanning representation, translation, alignment, fusion, co-learning, and missing data handling. Representation learning considers how to encode data from heterogeneous sources into mathematical structures suitable for computational processing. Approaches range from joint representations that project all modalities into one common space to coordinated representations that maintain modality-specific encodings linked through correlation constraints. Translation covers the process of converting information from one modality to another, as in image captioning or text-to-image synthesis. Alignment establishes correspondences between elements across modalities; example tasks include matching spoken words to lip movements or associating image regions with textual descriptions. Fusion refers to how information is combined from multiple

10.48047/jocaaa.2026.35.01.68

modalities to make predictions or decisions, with strategies ranging from simple concatenation to sophisticated attention-based integration that weighs the contributions of each modality relative to context [5].

Co-learning exploits knowledge from one modality to improve learning in another, especially useful when one modality is associated with much more available training data or when supervision signals are available only for certain modalities. Transfer learning methods allow models trained on rich multimodal datasets to transfer knowledge and boost performance on tasks where some modalities are limited or expensive to acquire. The problem of missing modalities involves realistic settings in which some expected inputs are not available during either training or inference, and it requires the robust performance of a model despite incomplete information. The techniques involve learning to reconstruct a missing modality given others, developing modality-agnostic representations that are naturally robust to partial inputs, and training different pathway specializations that degrade gracefully given incomplete inputs [5].

Medical imaging, clinical notes, laboratory results, genomic data, and physiological signals-even multimodal AI applications in healthcare point toward its transformative potential. In real-world applications, one finds a lot of autonomous vehicle systems in continuous processing of camera images, LiDAR point clouds, radar measurements, GPS localization, and inertial motion data. Multi-modal object detection and semantic segmentation for autonomous driving research has identified a number of basic challenges regarding sensor calibration to ensure spatial alignment across modalities, temporal synchronization accounting for different sampling rates of sensors, and robust strategies for sensor fusion maintaining performance despite sensor failures or adverse environmental conditions affecting only a few modalities. Deep learning architectures for autonomous driving rely on early fusion by concatenating raw sensor data, late fusion by combining the predictions from modality-specific models, or hybrid approaches that fuse at some intermediate levels in their feature computation. Natural language processing applications have been revolutionized through systems engaging in image-based conversations, generating descriptions, answering questions about visual content, and processing complex documents integrating text, figures, and tables.

Technical Challenge	Definition	Application Domain
Representation	Encoding heterogeneous data into mathematical structures	Cross-modal processing
Translation	Converting information between modalities	Image captioning, synthesis
Alignment	Establishing cross-modal correspondences	Speech-to-text matching
Fusion	Combining multi-modal information for decisions	Autonomous vehicles
Co-learning	Leveraging knowledge across modalities	Transfer learning tasks
Missing Data Handling	Maintaining performance with incomplete inputs	Real-world deployments

Table 2: Core Technical Challenges in Multimodal Machine Learning [5, 6]

IV. Edge Computing and Specialized Hardware Architectures

Advanced AI models have traditionally been computationally intensive and thus needed to be performed centrally in data centers. However, a number of factors are converging that are forcing a vital shift toward edge computing, where AI inferences happen right on the local devices themselves. This reimagining of architecture overcomes some crucial limitations in latency, privacy, energy efficiency, and network connectivity. Key benefits of implementing edge AI include reduced latency for real-time applications, preservation of privacy through local processing of data, lower network bandwidth use, and sustained operation in case of disrupted connectivity.

Edge intelligence is an intersection of AI and edge computing to offer intelligent processing at network edges instead of some cloud facilities. The whole framework contains a chain of three tiers: cloud computing for ample computational resources, often used for model training and complex analytics; edge computing, usually performed on a local server or gateway, for intermediate processing requiring moderate resources; and end devices, which include anything from smartphones to sensors, for lightweight inference. Such a hierarchical structure allows workload distribution to optimize among application requirements, computation constraints, and communication costs. Model partitioning strategies divide neural networks across tiers, executing initial layers on resource-constrained devices while offloading computationally intensive layers to edge servers or the cloud when needed. Early exit mechanisms enable models to terminate the inference at intermediate layers in the case of meeting confidence thresholds, therefore reducing computational requirements in straightforward cases while reserving the full model capacity for ambiguous input [7].

Training edge AI models presents unique challenges that require adaptation of classical approaches. Federated learning learns models collaboratively using distributed devices without collecting raw data centrally, hence solving privacy issues and enabling the advantages of diversity in data. Nonetheless, edge federated learning is prone to practical issues, including limited device computing power, sporadic connections, dissimilar distribution of participant data, and communication expenses that take over training time. On-device learning facilitates model adaptation to user-specific or local environmental conditions by adopting continual learning methods that update models incrementally using newly observed data. Transfer learning reduces the training burden by adapting larger pre-trained models to specific contexts of edge deployment by fine-tuning them on limited local data. Model compression techniques include knowledge distillation, quantization, pruning, and architecture search, which yield compact models for resource-constrained devices at adequate accuracy levels [7].

Neuromorphic computing represents a brain-inspired paradigm with revolutionary potential for ultraefficient edge AI through event-driven processing inspired by biological neural mechanisms. Recent advances in neuromorphic engineering demonstrate integration of algorithmic innovations, neuromorphic processors, and intelligent systems that support real-time processing at minimal energy consumption. These neuromorphic processors use the principle of spiking neural networks, wherein the artificial neurons communicate discrete events instead of continuous activations, hence radically reducing energy consumption for temporal pattern recognition and sensory processing tasks. The hardware implementations of digital, analog, and mixed-signal systems have to trade off between flexibility, efficiency, and biological fidelity. Digital neuromorphic systems, while being programmable and precise, consume more power; on the other hand, analog implementations achieve high energy efficiency at the expense of reduced configurability and susceptibility to device variations. Applications spanning robotics, always-on monitoring, brain-machine interfaces, and adaptive control demonstrate neuromorphic computing's potential for edge intelligence requiring continuous, low-power perception and decisionmaking in dynamic environments.

Architectural Tier	Computational Resources	Primary Functions	Processing Type
Cloud Computing	Massive resources	Model training, complex analytics	Centralized
Edge Computing	Moderate resources	Intermediate processing	Distributed
End Devices	Limited resources	Lightweight inference	Localized
Model Partitioning	Variable allocation	Layer distribution across tiers	Hierarchical
Neuromorphic Processors	Ultra-efficient	Event-driven processing	Brain-inspired

Table 3: Edge Computing Architecture Tiers and Processing Characteristics [7, 8]

V. Privacy-Preserving and Transparent AI Techniques

The first is preserving privacy during model training and deployment while offering transparency and interpretability of decision-making processes, especially as AI systems process increasingly sensitive data and make progressively consequential decisions. Complementary approaches such as federated learning and explainable AI address these fundamental concerns affecting the adoption of AI within regulated industries and high-stakes applications. The techniques herein respond to a growing recognition that AI trustworthiness depends not just on predictive accuracy but also on verifiable privacy protection and comprehensible reasoning processes accessible to stakeholders with varied technical expertise.

Federated learning basically recasts the model training paradigm of machine learning for data privacy by enabling model training across distributed data sources without ever exposing raw data from its original location. Accordingly, the training process happens in iterative rounds where model instances are locally trained by edge devices or institutional holders of private data, followed by the uploading of only model updates to a central coordinator that sums up distributed contributions to generate a better global model. This architectural position allows for strong privacy advantages due to the fact that data remains at all times within organizational and jurisdictional boundaries during training, thereby meeting data localization requirements while minimizing exposure to unauthorized access. The general federation learning structure includes horizontal federated learning, in which participants share the feature space but have distinct sample populations; vertical federated learning, in which participants share samples but have dissimilar features; and federated transfer learning, which deals with the situation where the participants share both samples and features.

Technical issues: Federated learning compares to traditional centralized training in the following critical technical issues. Statistical heterogeneity across participants-the data is non-identically distributed, classes can be imbalanced, or different feature characteristics-renders the model convergence poor and worsens the final performance compared to training on combined data. Communication efficiency becomes critical when the training involves a large number of participants with low bandwidth or intermittent connectivity, which demands gradient compression mechanisms, sparse update mechanisms, and optimized communication protocols that can minimize the data that needs to be transmitted. System heterogeneity, where participants employ quite heterogeneous hardware with diverse computational capability and storage constraints, raises a challenge of careful orchestration to prevent bottlenecks from slow participants while

10.48047/jocaaa.2026.35.01.68

continuing to maintain contribution fairness. Security: Security considerations go toward threats such as data poisoning attacks, in which malicious participants corrupt the training via adversarial examples, model poisoning by submitting malicious gradients, and inference attacks aimed at eliciting information about the training data from model updates [9].

Healthcare applications epitomize this transformative potential of federated learning by collaborative model development without necessarily sharing patient records. Various research papers have demonstrated that federated learning enables the creation of diagnostic models utilizing relevant clinical data in a distributed fashion without breaching the strict privacy protections inherent in medical ethics and regulations. Swarm learning extends federated principles through blockchain-based coordination that eliminates the use of centralized servers, hence allowing peer-to-peer collaboration among institutions. Applications across COVID-19 detection, cancer classification, and treatment response prediction show that decentralized training achieves comparable performance to centralized training methods while providing strong privacy guarantees. The swarm learning framework utilizes blockchain for transparent coordination, cryptographic techniques that secure model updates during transport, and consensus mechanisms that make sure participants agree on global model states. Financial institutions are using federated learning similarly: fraud detection models benefit from broader transaction pattern visibility across institutions without exposing proprietary customer data or violating financial privacy regulations. Explainable AI addresses complementary transparency challenges as models become increasingly complex and autonomous. Modern deep learning architectures use millions or billions of parameters in complex structures, making internal decision-making processes effectively impossible to interpret by conventional inspection. However, there is a wide range of uses that demand knowledge of why systems have arrived at specific conclusions, to either check the suitability of the reasoning or to provide regulation or ethical responsibility. The informational requirements of transparency vary with the stakeholder concerned: end users require the ability to make decisions based on intuitive explanations, domain experts require the ability to validate either based on clinical or technical suitability, regulators require the ability to verify compliance by relying on detailed documentation, and system developers require the ability to debug to identify failure modes. There must be a compromise between explainability methods to achieve fidelity that ensures model behavior is well represented, interpretability that the method is understandable by humans, and completeness that ensures that all factors are considered in the decision.

Technique/Challenge	Category	Key Feature	Impact Domain
Horizontal Federated Learning	Framework Type	Shared feature space	Cross-organizational training
Vertical Federated Learning	Framework Type	Shared samples	Multi-feature integration
Statistical Heterogeneity	Technical Challenge	Non-identical data distribution	Model convergence
Communication Efficiency	Technical Challenge	Bandwidth optimization	Training scalability
Swarm Learning	Implementation	Blockchain coordination	Healthcare collaboration

10.48047/jocaaa.2026.35.01.68

System Heterogeneity	Technical Challenge	Diverse hardware capabilities	Orchestration complexity
----------------------	---------------------	-------------------------------	--------------------------

Table 4: Privacy-Preserving AI Techniques and Implementation Challenges [9, 10]

VI. Enterprise Integration and Governance Frameworks

This active engagement of AI functionality into the current systems presents a decisive stage in the maturation of technology, where the emphasis is now moving beyond individual applications to the seamless addition of mature enterprise business processes. With invisible AI, intelligent capabilities are deeply embedded in applications and platforms that professionals work with today; it increases productivity for various organizational functions while reducing adoption friction. At the same time, the increased scope of AI deployment has moved governance concerns from a mere technical implementation detail to one of strategic organizational and societal priorities, which require a comprehensive framework of accountability, fairness, transparency, and compliance along the AI lifecycle.

Enterprise AI integration is more likely to occur in a cluster of functional areas where the intelligent functionality supplements, instead of substitutes, existing business processes. The software development environment features code completion, bug reports, and automated testing built on AI that maintain shorter development cycles and lower errors. The use of AI in marketing platforms to generate content for audience segmentation, campaign optimization, and performance analytics can be used to engage customers more specifically. Systems of customer relationship management use AI to score customer leads, churn, and customized communication plans. This integration approach has a number of benefits over single deployments, such as the use of familiar interfaces that require less training, the use of existing data sources, and authorization models that have reduced security management and can deliver incremental value through gradual improvement as opposed to a total overhaul of the infrastructure. Collaborative intelligence frameworks stress human-AI collaboration where machines and humans enhance the powers of each other, rather than AI automating human tasks simply.

Research has distinguished three waves of business transformation: standardized processing, where machines execute routine tasks more efficiently and predictably than individual humans; expert-like decision-making, where systems use data-driven insights to make decisions that no human could; and collaborative intelligence, where humans and AI interactively enhance each other's powers. Of these, the third is the most revolutionary: machines greatly extend the cognitive powers of humans, but humans ensure the responsible, explainable operation of AI systems within ethical bounds. Humans train the systems in the first place by teaching and explaining the nuances; humans interact with machines to harness their power for heavy analysis of complex situations; and humans create sustainable use of AI by monitoring ethical operation and detecting problematic behaviors of the systems. This collaboration effectively requires reimagination of business processes around human-AI interaction rather than just inserting AI into the workflows. It is the place where humans excel in leadership, teamwork, creativity, and socialization, and machines excel in speed, scalability, quantitative abilities, and endless repetition. Effective implementations capitalize on such complementary advantages using hybrid modes where AI is used to perform data-intensive pattern recognition, and humans are used to offer contextualization, ethical checkpoints, and adjustment to new circumstances. It requires them to invest in new capabilities, including coming to terms with the capabilities and limitations of the AI systems, the management of the workflow with AI-based enhancements, and the promotion of human oversight in the development of autonomous systems. Governance frameworks offer frameworks that are required to responsibly deploy AI, and to deal with

10.48047/jocaaa.2026.35.01.68

accountability, fairness, and transparency of compliance in the course of development and operational stages.

Comprehensive AI governance covers not just clear accountability structures that define roles and responsibilities for system oversight, risk assessment frameworks that systematically evaluate potential negative impacts, but also fairness and bias mitigation that ensures equity in treatment across demographic groups, and transparency requirements that define what information about system design and operation is disclosed. Algorithmic transparency for smart cities can be taken as the exemplar for governance challenges in public sector contexts where AI systems make decisions with implications for citizen welfare, resource allocation, and service delivery. There are numerous variants of smart city implementations with the use of AI in traffic control, power supply, civic safety, and municipal services. The issue of sufficient amounts of transparency is being raised with the balance towards operational security and competitive concerns. The transparency requirements should accommodate the needs of a large scope of the various stakeholders, either due to technical limitations or proprietary interests. The full algorithmic disclosure, i.e., the full disclosure of training data, model structures, and decision processes, is usually not practical, because it may be insecure or simply too technical to be interpreted by the majority of stakeholders.

Other options include functional transparency-explaining what goes into and comes out of the system and how it generally works without disclosing specific implementation details-and audit-based transparency, which provides a way for independent experts to assess systems under controlled conditions that verify compliance without exposing sensitive elements to the public. More recent regulatory frameworks are increasingly demanding systems, especially those that may impact individual rights. New legislation is taking up such issues as algorithmic accountability, automated decisions, high-risk AI applications, and sector-specific requirements [12]. Organizations will need to monitor changing regulatory requirements, establish mechanisms for compliance, and maintain records that demonstrate adherence to relevant standards, as governance frameworks change to balance innovation enablement with safeguards appropriate for trust and ensuring AI serves human interests.

Conclusion

The Central African Republic is reportedly moving closer to becoming a disaster for everyone. Artificial intelligence represents a basic paradigm shift in computation, evolving from solitary task execution toward integrated systems with autonomous reasoning, multiple perception modalities, and collaborative mechanisms that preserve privacy. Chain-of-thought prompting and foundation models have established much higher levels of complex reasoning, emerging from architectural scaling, but also autonomous agents indicating proactive problem-solving by means of multi-step planning and coordination. Multimodal architectures synthesize a variety of data streams for complete contextual understanding necessary for complex real-world applications. Edge computing and neuromorphic processors address computational efficiency through the use of localized processing and brain-inspired architectures, respectively, and can be safely deployed on resource-constrained devices. Federated learning frameworks preserve privacy by decentralized training; explainable AI techniques facilitate regulatory accountability and build user trust. Integration into enterprises via collaborative intelligence frameworks positions AI as complementary to human capabilities, rather than a replacement technology. Comprehensive governance structures addressing accountability, fairness, and transparency provide bases for responsible deployment across sensitive domains. These converging technological advances set up AI systems with the capability to transform healthcare diagnostics, independent transportation, financial services, and urban infrastructure, and are in a position to remain aligned with human values and societal interest via robust oversight mechanisms.

References

- [1] Jason Wei et al., "Chain-of-thought prompting elicits reasoning in large language models," arXiv:2201.11903v6, 2023. Available: <https://arxiv.org/pdf/2201.11903>
- [2] Rishi Bommasani et al., "On the opportunities and risks of foundation models," *arXiv:2108.07258, 2022. Available: <https://arxiv.org/abs/2108.07258>
- [3] Lei Wang et al., "A survey on large language model-based autonomous agents," Springer, 2024. Available: <https://link.springer.com/article/10.1007/s11704-024-40231-1>
- [4] Natasha Jaques et al., "Social influence as intrinsic motivation for multi-agent deep reinforcement learning," Proceedings of Machine Learning Research, 2019. Available: <https://proceedings.mlr.press/v97/jaques19a.html>
- [5] Paul Pu Liang, Amir Zadeh, and Louis-Philippe Morency, "Foundations and recent trends in multimodal machine learning: Principles, challenges, and open questions," arXiv:2209.03430, 2023. Available: <https://arxiv.org/abs/2209.03430>
- [6] Di Feng et al., "Deep Multi-modal Object Detection and Semantic Segmentation for Autonomous Driving: Datasets, Methods, and Challenges," arXiv:1902.07830, 2020. Available: <https://arxiv.org/abs/1902.07830>
- [7] Zhi Zhou et al., "Edge Intelligence: Paving the Last Mile of Artificial Intelligence with Edge Computing," arXiv:1905.10083v1, 2019. Available: <https://arxiv.org/pdf/1905.10083>
- [8] Mike Davies et al., "Advancing Neuromorphic Computing With Loihi: A Survey of Results and Outlook," Proceedings of the IEEE, Volume 109, Issue 5, 2021. Available: <https://ieeexplore.ieee.org/document/9395703>
- [9] Qiang Yang et al., "Federated machine learning: Concept and applications," ACM Transactions on Intelligent Systems and Technology (TIST), Volume 10, Issue 2, 2019. Available: <https://dl.acm.org/doi/10.1145/3298981>
- [10] Stefanie Wernat-Herresthal et al., "Swarm learning for decentralized and confidential clinical machine learning," Nature, Volume 594, Pages 265–270, 2021. Available: <https://www.nature.com/articles/s41586-021-03583-3>
- [11] H. James Wilson and Paul R. Daugherty, "Collaborative intelligence: Humans and AI are joining forces," Harvard Business Review, 2018. Available: <https://hbr.org/2018/07/collaborative-intelligencehumans-and-ai-are-joining-forces>
- [12] Robert Brauneis & Ellen P. Goodman, "Algorithmic transparency for the smart city," Yale Journal of Law and Technology. Available: <https://yjolt.org/algorithmic-transparency-smart-city>