

Toward Explainable and Trustworthy Federated Big Data Intelligence in Secure Edge-Cloud Ecosystems: A Conceptual Framework and Research Agenda

Khaled Mohamed Mohamed Khalifa

Email: khalid.mohamed.khalifa@gmail.com

Abstract: Edge-cloud big data ecosystems are changing many areas, such as healthcare, smart cities, manufacturing IoT settings, cybersecurity, self-driving systems, mobile platforms, and financial analytics. Additionally, these groups produce a lot of data that is dispersed, unique, changes based on the situation, is sensitive to delays, and is usually private. Classical centralized cloud analytics are having more and more issues with bandwidth, latency, single points of failure, and keeping private data and power to make decisions in one place. Multiple edge clients can work on the same model at the same time with federated learning, and the raw data stays close by. Even so, federated learning doesn't guarantee accurate knowledge by itself. Although model changes can keep private data safe, hacked clients can ruin global learning, non-IID data can lead to performance gaps, and choices made by many people may still be hard to explain, audit, or control. This article suggests a trustworthy way to think about shared big data information that can be communicated in secure edge-cloud settings. Additionally, it includes studies on edge-cloud computing, shared learning, AI that can be explained, privacy-preserving analytics, security, and AI control. Additionally, the paper includes a study plan, an assessment tool, a trust-risk mapping, evaluation factors, and governance development standards. When making a case, the main point is that trustworthy federated intelligence is more than just distributed learning that protects privacy. It's a sociotechnical paradigm that includes privacy protection, security robustness, explanation quality, accountability, auditability, fairness, the ability to be deployed, and meaningful human oversight.

Keywords: Federated learning; Explainable artificial intelligence; Trustworthy AI; Edge-cloud computing; Big data intelligence; Privacy-preserving analytics; Secure distributed learning; AI governance

1. Introduction

With the rise of edge-cloud systems, the way that current intelligence is put together has changed. More and more data is being stored close to people, devices, monitors, tools, and work areas, rather than in large data centers. This takes place in smart towns, financial data, self-driving cars, mobile services, industrial Internet of Things (IIoT) systems, and healthcare tracking. Some things have changed, but not just the way things work. It's a sign of a bigger shift in big data intelligence: data are big in terms of volume, speed, and variety, but they're also spread out geographically, statistically diverse, private, and deeply rooted in local settings. Edge devices can keep an eye on what users do, how the environment changes, health trends, signs in the industry, or activities that could be risks. When making a case, the main point is that trustworthy federated intelligence is more than just distributed learning that protects privacy. It's a sociotechnical paradigm that includes privacy protection, security robustness, explanation quality, accountability, auditability, fairness, the ability to be deployed, and meaningful human oversight. In these cases, intelligence is more than just putting data in one place and using cloud-based tools to figure out what it all



means. For privacy, security, delay, and control rules to be followed, local edge nodes and cloud-level systems must work together (Abbas et al., 2018; Mao et al., 2017; Roman et al., 2018; Shi et al., 2016; Satyanarayanan, 2017).

For many uses of big data, centralized cloud analytics has been helpful. However, its flaws are becoming more apparent in places where privacy is important and resources are limited. When you store and move raw data, it's possible that private data like medical, financial, operational, or security data will be made public. A lot of high-frequency data streams coming from edge devices can also make the bandwidth tight. When you depend on cloud technology that is far away, it could slow down things that need to work quickly, like automatic systems, tracking job safety, finding hackers, and emergency healthcare. Centralized designs can also make single points of failure worse, share the power to make decisions, and make it harder to follow data security rules that say data can't move across countries or needs to be processed in a certain way. Because of these worries, people have created distributed intelligence models where data sources are nearer to where analytics and learning take place (Abbas et al., 2018; Mao et al., 2017; Shi et al., 2016).

Federated learning is one of the most important ways to make sure that global machine learning takes privacy into account. In a shared space, clients train models on their own data. They don't send raw records; instead, they share changes to the model with a server that coordinates or an aggregate way. Edge-cloud ecosystems like this set-up a lot because it keeps local data ownership, reduces the amount of raw data that needs to be moved, and allows institutions, devices, or domains that can't easily share data to work together (Aledhari et al., 2020; Kairouz et al., 2021; Lim et al., 2020; McMahan et al., 2017; Nguyen et al., 2021; Yang et al., 2019; Zhang et al., 2021). For federated big data intelligence to work, data from different sources is used to build a group model while keeping local data separate. This is called federated learning.

The main idea behind this work is that reliable edge-cloud intelligence does need joint learning, but it doesn't just need it. For federated learning to work, raw data can't be directly controlled. However, this doesn't protect privacy, security, fairness, auditability, responsibility, or human control. Grain and model parameters can be used in member inference, gradient reconstruction, or feature reconstruction attacks to get to local data (Nasr et al., 2019; Zhu et al., 2019). Foul play, hidden attacks, strange behavior, and clients that have been hacked are some other dangers that shared systems may face. This is very important in big edge cases where you can't trust everyone (Bagdasaryan et al., 2020; Blanchard et al., 2017; Mothukuri et al., 2021). Lack of IID data can also cause convergence to be unstable, model quality to be uneven, and customers to have hidden bias (Kairouz et al., 2021; Li et al., 2020).

Things that are easy to explain are harder to accept. AI methods that can be described have been the subject of a lot of study. Asdai and Berrada (2018), Barredo Arrieta et al. (2020), Guidotti et al. (2018), Gunning and Aha (2019), and Ribeiro et al. (2016) list some of these: local- and global-explanations, feature attribution, alternative reasoning, and interpretable substitute models. Some decisions may need to be made privately in linked edge-cloud systems, then compared across clients, gathered at higher levels, and then shared without showing private client data. People in different areas may have very different views because they see different types of people, work situations, or dangers. When global reasons make sense, they can hide the behaviors of minority clients or problems that are unique to the situation. Because they show private traits, trends, or links that only one client can see, the results of explanations can also be a privacy route. Explainability shouldn't be added as a feedback layer after the fact in shared systems; it should be put in during the global learning process.

There is a lot of value in both governance and duty. A trustworthy shared edge-cloud intelligence system should let you keep track of reports, paperwork, risk rating, job sharing, compliance tracking, human review, and the ability to challenge it. There may not always be clear who is to blame when there are many people involved in a system: a global model fails, a client makes changes that are unfair or bad, the reasons given are wrong, or a shared choice has a bad effect on a user. Good spread intelligence requires governance, not just more management. Expert Group on Artificial Intelligence (European Commission, 2019; Floridi et al., 2018; Jobin et al., 2019; Mittelstadt, 2019; National Institute of Standards and Technology [NIST], 2023) says that rules only mean something if they are put into action through institutional processes, accountability mechanisms, and operational controls. We agree with this point of view.

2. Methodological Foundations

2.1 Edge-Cloud Big Data Ecosystems

There are edge devices, middle edge servers, and centralized or semi-centralized cloud platforms in edge-cloud big data ecosystems. These all work together to help make data, process it locally, draw conclusions from models, and organize at the cloud level. Edge nodes are things like smartphones, sensors, ports, industrial computers, cars that are linked together, medical devices, tracking systems, and security tools. Every so often, these nodes send out data that shows things like acts happening in the neighborhood, user behavior, weather conditions, machine states, and rare events. The data is full of context. For standard centralized analytics, raw data is sent to the cloud to be kept and worked on. In edge-cloud analytics, work is split up and done on many levels. The processing is moved closer to the source of the data, but cloud tools are kept for storage, organization, and large-scale analysis (Abbas et al., 2018; Mao et al., 2017; Roman et al., 2018).

The big data in the edge cloud are not only big, but they are also spread out, different, changing, and often private. Because of these things, this change in design is important. Smart health care systems may collect body marks, IoT systems may send out unique production signs, cybersecurity platforms may look at network logs, and mobile services may look at patterns of where people are or what they're doing. In this case, you can't just trust the forecast tool. It needs to be put into tracking, data management, device stability, communication security, and the strength of infrastructure as well. Edge-cloud intelligence needs ways to stay independent at the local level while letting everyone learn together and make smart decisions (Shi et al., 2016; Satyanarayanan, 2017).

2.2 Federated Learning and Federated Big Data Intelligence

Federated learning is a way to set up machine learning that is spread out in places that care a lot about data security. To train a local model, each client uses its own data. It then tells a central server or an aggregation method about any changes it makes to the model. The boss is putting these changes together to make a world plan. After that, this model can be sent to clients over and over again through different contact rounds (Aledhari et al., 2020; Kairouz et al., 2021; McMahan et al., 2017; Yang et al., 2019; Zhang et al., 2021). Collaborative learning lets people from all over the world help build models without having to send a lot of raw data.

This idea goes further than just making models work better with Federated Big Data Intelligence. People and technology are set up in a more general way so that clients can get help with analytics, risk tracking, and making choices without having to put all of their raw data in one place. Clients in an edge-cloud setting might have different working goals, network connections, computer power, data quality, data volume, and situations in their own areas. These sites can learn from each other even when they can't or shouldn't share data directly (Lim et al., 2020; Nguyen et al., 2021; Rieke et al., 2020; Xu et al., 2021). One thing that should not be thought of as a full trust system is joint learning. It should be seen as a technology base instead. It lowers one kind of privacy risk, but it doesn't fix issues with bias, leaks, responsibility, or privacy by itself (Mothukuri et al., 2021; Zhu et al., 2019).

2.3 Explainable Artificial Intelligence in Distributed Environments

Explainable AI wants to make machine learning systems easy to understand and that can be held accountable by the right people. The XAI methods can explain results in a very specific way or in a more general way by showing how the model works in general. They could also use rule extraction, alternative reasoning, feature assignment, replacement models, or analysis based on examples. Some important papers that have been written recently are by Adadi & Berrada (2018), Barredo Arrieta et al. (2020), Guidotti et al. (2018), Gunning & Aha (2019), and Ribeiro et al. (2016). They all talk about things like how right, stable, interpretable, valuable, and in line with subject knowledge answers are. Because a lot is at stake, answers are more than just facts. Expert proof, regulatory review, trust tuning, and making smart decisions are all things they help with.

In scattered edge-cloud settings, it's harder to explain how models work because each place has its own working conditions and data trends. Two clients may not understand each other's ways of thinking. However, a world report

might only show how the model usually works and not local risks, harms that are unique to the place, or patterns of behavior among group clients. With pooled explainability, it is important to compare, plan, protect, and check both local and global reasons. In terms of privacy, you should say that even if the raw records stay in the same place, the results can still show private trends. It is important for a reliable shared XAI system to keep an eye on explanation drift and make sure that sharing explanations is safe. Also, answers should be able to be used by technology experts, subject workers, managers, lawmakers, and users who are affected.

2.4 Trustworthy AI and Governance in Edge-Cloud Systems

It's not enough for AI to be correct in predictions for it to be accepted. It has to be open, safe, strong, private, fair, responsible, able to be inspected, and controlled by real people. These needs are stronger in edge-cloud groups where data, technology, models, and ways of making decisions are spread out among many people and places. Individuals must have faith in stable local clients, secure communication channels, risk-free model changes, comprehensive explanations, the discovery of harm, and the ability to hold individuals responsible when things go wrong.

Legitimate AI ideas are put into action by government. Some of the things that make up this process are documentation, risk assessment, audit records, action plans, human review, keeping track of compliance, and giving duties. Sometimes, not a single person or group can see all the data, models, and decision settings in a shared edge-cloud system. This makes governance very important. So, for a system to be responsible, it needs both social duty and technology defenses. This is in line with other research on AI governance that says privacy, security, fairness, explainability, and auditability should be seen as unified needs rather than separate moral goals (European Commission High-Level Expert Group on Artificial Intelligence, 2019; Floridi et al., 2018; Jobin et al., 2019; Mittelstadt, 2019; NIST, 2023).

Table 1. Conceptual distinction between federated learning, explainable AI, and trustworthy federated edge-cloud intelligence.

Dimension	Federated learning	Explainable AI	Trustworthy federated edge-cloud intelligence
Primary aim	Distributed model training	Model transparency and interpretation	Secure, explainable, auditable, and accountable distributed intelligence
Data assumption	Raw data remain local	Often assumes centralized access to model behavior or data	Raw data remain local and explanation exchange is privacy-aware
Main output	Global or personalized model	Explanation, attribution, or interpretation	Governed intelligence process supported by evidence and oversight
Key risk	Update leakage, poisoning, non-IID drift	Misleading, unstable, or unusable explanations	Combined privacy, security, interpretability, fairness, and accountability failure
Evaluation focus	Accuracy, convergence, communication efficiency	Fidelity, stability, interpretability, usefulness	Trust, privacy, security, explainability, fairness, governance, latency, and deployment feasibility
Limitation	Not automatically interpretable or governable	Not automatically private or secure	Requires integrated socio-technical system design

3. Trust Challenges in Federated Edge-Cloud Big Data Intelligence

3.1 Data Heterogeneity and Client Drift

An important issue with sharing edge-cloud knowledge is that the clients are not all the same. Individual clients' data may be different in where it comes from, how good it is, how many of them there are, how stable it is over time, and what it means in each place. In centralized learning, data are shared before training, but this is not the same thing. A lot of the time, data in an edge cloud is spread out in the same way and is not kept separate (non-IID). Only a small part of the larger data area may be visible in places like a hospital, office, cell phone, or security point. That is why local training could lead to ways to get better that are specific to the client and not the general goal.

Ethic and technology change when customers switch. In addition, it can make global models that don't work the same way for all clients and make convergence and consolidation less safe. If a client is small or not well-represented, they might not be shown enough, while clients who are big or well-represented may have a bigger effect on world learning. That's why variety is bad for both trust and the business. Some customers may get different results and reasons, and if no one is watching, the system might be biased, unfair, or harmful in ways that can't be seen (Kairouz et al., 2021; Li et al., 2020).

3.2 Privacy Leakage and Confidentiality Risks

As a result, group learning is sometimes called "privacy-preserving." Although this is a good explanation, it's not all there is to say. Concerning model changes, slopes, factors, information, chat trends, and answer results, there is still the same privacy risk. Sharing updates can be used by gradient leaks and membership inference attacks to uncover whether a record was in training or to get back parts of local training data (Nasr et al., 2019; Zhu et al., 2019). The dangers are even greater for edge-cloud systems that handle private, secure, or medical data.

Additionally, giving reasons can cause leaks. Unusual traits, sensitive behaviors, or signs that are unique to the situation could affect a client's choice. In situations where there aren't many customers or where the groups aren't all the same, combining theories may reveal minority trends. For trustworthy shared intelligence to work, privacy must be protected during training, data collection, explanation creation, explanation sharing, and audits. You could use secure aggregation, differential privacy, encryption, and secure multiparty computing, but these aren't separate security measures. Instead, they're ways to improve the system (Bonawitz et al., 2017; Dwork & Roth, 2014; Mothukuri et al., 2021).

3.3 Security Threats and Malicious Client Behavior

In shared edge-cloud systems, bad behavior is possible because learning needs information from many clients. Some clients may be broken into, insecure, or bad on purpose when they are in an open or semi-open setting. Threats that spread poison can make it hard to gather information and teach the global model bad habits. As Bagdasaryan et al. (2020) say, backdoor attacks can change what happens when certain things happen while keeping normal behavior on common inputs. Blonchard et al. (2017) say that byzantine clients may send changes that are random, fake, or meant to make learning harder.

When used with edge clouds, these threats are worse because edge nodes might not have as much real security, smaller computer resources, links that don't work as well, and different security settings. Someone who has got into a workplace monitor, mobile client, or local port can use it to change the whole learning process. Because of this, strong aggregation, trust scores, update validation, and ways to deal with clients that look sketchy should all be part of shared coordination. Federated learning could keep data local and still give us bad or dangerous intelligence if these safeguards aren't in place (Mothukuri et al., 2021; Roman et al., 2018).

3.4 Explainability Gaps in Federated Systems

Using current XAI methods on the final global model alone won't be enough to explain federated systems. Explainability gaps happen when local models learn from different types of data and situations. One client may not

care about or need a certain trait, while another may find it useful or damaging. If you give a world answer, it might hide these differences by giving a stable-looking general meaning that doesn't hold true in certain places.

Not-IID data make figuring out what an answer means even harder. There may be real differences between local explanations, but there may also be signs of drift, bias, bad data, hostile manipulation, or an unstable explanation method. To be sure that shared XAI is reliable, it needs to be checked for security, data leaks, correct explanations, and usefulness to stakeholders. Additionally, there should be a government that keeps track of how explanations were made, who can see them, and how differences between local and global explanations are resolved (Adadi & Berrada, 2018; Barredo Arrieta et al., 2020).

3.5 Accountability, Auditability, and Governance Challenges

Sharing edge-cloud information hasn't had a lot of work done on accountability. That one company might be in charge of the data, models, marketing, and keeping track of AI when it is handled. Of shared systems, various individuals are in charge of various parts. These groups include people who own the data, clients at the edges, aggregation services, cloud providers, infrastructure operators, security teams, model writers, and decision-makers further down the line. Sharing a model might not be the best option because of inaccurate local data, changes meant to hurt, poor collection, inadequate privacy controls, bad reasoning, bad management, or a combination of these issues.

Allowing audits is very important for this reason. It's safe to assume that shared intelligence will keep track of all the exchanges that happen between clients, updates, validation rules, model versions, outputs, privacy settings, security alerts, and people. Risk logs should keep track of more than just technical errors. Governance issues like repeated "anomaly flags," clients being left out without a reason, differences between local and global answers, and gaps in who is responsible for what should also be recorded. A person should always be in charge of decisions that are made automatically that affect safety, rights, resources, or legal chores. Based on the European Commission High-Level Expert Group on AI (2019) and NIST (2023), a common system could be very advanced but not secure in terms of its institutions if these steps are not taken.

Table 2. Trust challenges, risks, and conceptual requirements in federated edge-cloud intelligence.

Trust challenge	Main risk	Why it matters in edge-cloud systems	Required conceptual response
Non-IID data	Biased or unstable global model	Clients observe different populations, devices, and environments	Monitor local-global performance and explanation divergence
Privacy leakage	Exposure through updates, metadata, or explanations	Raw-data locality does not eliminate inference risk	Use secure aggregation, differential privacy, and privacy-aware explanation sharing
Poisoning and backdoors	Corrupted global model or targeted unsafe behavior	Edge clients may be compromised or weakly protected	Apply robust aggregation, anomaly detection, and trust scoring
Weak explainability	Opaque or misleading distributed decisions	Operators and domain experts need understandable evidence	Provide local explanations, global explanations, and alignment checks
Poor governance	Unclear responsibility and weak auditability	Multiple actors participate in data, models, infrastructure, and decisions	Create audit trails, accountability allocation, documentation, and human oversight

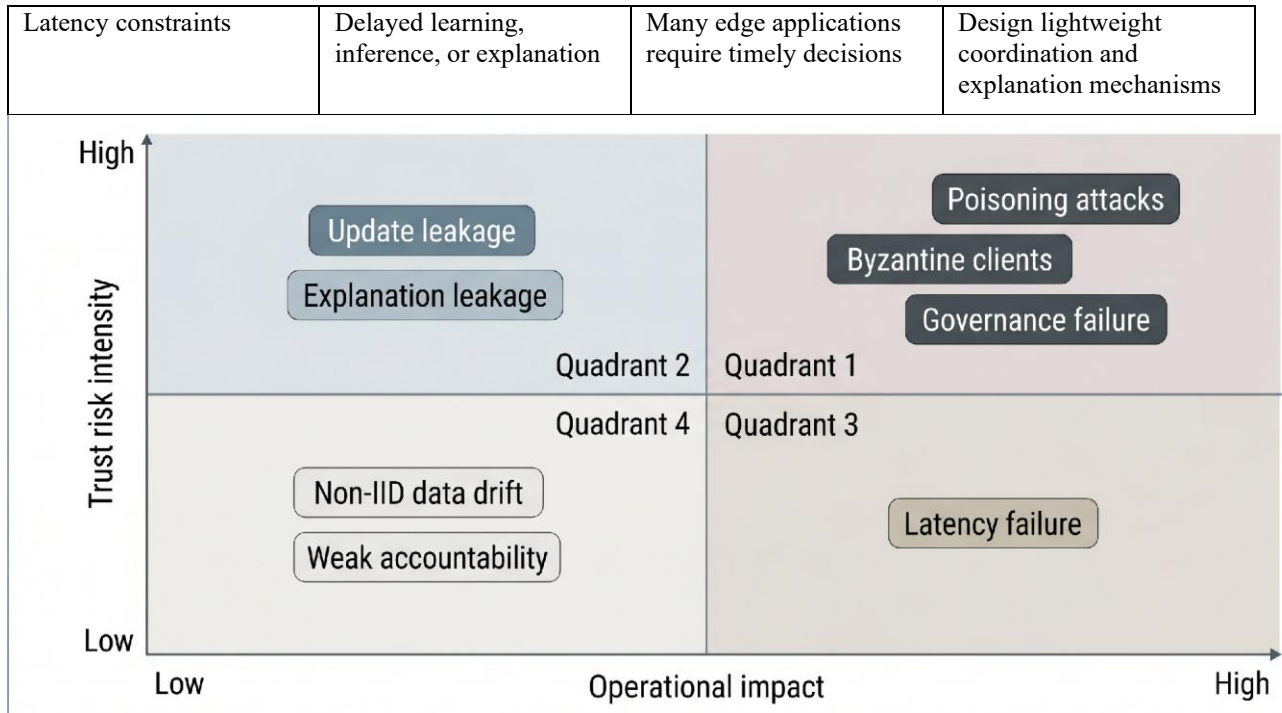


Figure 1. Trust-risk map for federated edge-cloud big data intelligence.

The most dangerous dangers are not just those that make the model less accurate, as shown in Figure 1. In the high-risk/high-impact region, you'll find poisoning attacks, Byzantine clients, and governance failure, all of which can hurt both IT performance and trust in institutions. The idea that local data keeping is enough is called into question by update and justification leaks, which are high-intensity confidence risks. Lack of accountability and non-IID movement may add up over time, leading to secret bias, shaky reasons, or confusing obligations. Because many edge-cloud apps need quick and light trust methods, latency failure is put in an area with a lot of significance.

4. Proposed Analytical Taxonomy

4.1 Taxonomy Rationale

There are still a lot of different ideas in the writings on shared edge-cloud intelligence. A lot of study on federated learning is done on global optimization, communication efficiency, aggregation, and learning that doesn't use an ID. Model openness, credit, interpretability, and human knowledge are some of the things that explainable AI study looks at. Leakage, encryption, safe aggregation, differential privacy, poisoning, and Byzantine resilience are all technical issues that are often looked at separately in privacy and security studies. Governance study is mostly about issues like duty, transparency, paperwork, control, inspection, and the distribution of power. It is hard to understand shared edge-cloud intelligence as a single sociotechnical system because of this split.

The suggested classification sorts reliable shared edge-cloud intelligence into four groups that are linked to each other: learning design, explainability, privacy and security, and deployment and control. These factors don't work on their own. Learning architecture changes privacy risk and accountability limits; explainability changes openness but could lead to leaks; privacy controls can cut down on details in explanations; security monitoring could make governance more difficult; and deployment limits decide if trust mechanisms can work.

4.2 Learning Architecture Dimension

The learning design factor looks at how intelligence is set up at the client, edge, and cloud levels. Centralized cloud learning sends all of the raw data to one place so that it can be used for training and analysis. It may be possible

to do a lot of work with this, but it also makes people worry about privacy, speed, control, and energy. Edge-assisted learning moves some processes closer to the devices. This makes them faster, but they still need to work together. Federated learning lets users train in their own areas and share new models instead of raw data (McMahan et al., 2017; Kairouz et al., 2021). It depends less on a single central manager with peer-to-peer federated learning and more on edge or area servers for hierarchical federated learning.

When you accept each plan, different things happen. Centralized systems concentrate more risk in one area; shared systems allow more people to learn, but depend on clients being trustworthy; hierarchical systems may help businesses grow, but they make it harder to assign work; and peer-to-peer systems reduce reliance on the center while making it harder for teams to work together. Because of this, trustworthy building design should think about more than just how well it works. It should also think about duty, privacy, safety, and getting a lot of different clients.

4.3 Explainability Dimension

How well can technology experts, topic experts, system operators, lawmakers, and affected users understand how a distributed model behaves? For more information, see the explainability dimension. You can find out why a model got a certain result with a certain client by looking at the local description. There is a general explanation that tells you how the models work in the shared system as a whole. The local-global alignment test seeks to determine whether global reasons accurately describe how clients behave or whether they conceal nearby dangers. When you use human-centered explanation, the results of your explanation are checked to make sure they are clear, useful, and in line with what important people need.

Having power over answers is just as important as making them when there are different choices. It is best to write down explanations, keep them safe, check them, and give them to the right people in a way that lets them look them over and disagree with them. The use of XAI might not be a real part of shared intelligence if there isn't clear control.

4.4 Privacy and Security Dimension

It's important that data, model changes, details, records, and the way the system works don't get out or are changed by people who don't belong there. The privacy and safety part is all about this. Logic sharing that is aware of privacy, secure aggregation, encryption, secure multiparty computing, and differential privacy are some of the ways that privacy can be kept safe. It is safe to collect client updates with secure aggregation, and controlled noise is used for differential privacy to keep client-specific data from getting out (Bonawitz et al., 2017; Dwork & Roth, 2014).

Anomaly detection, trust score, client validation, strong aggregation, and poisoning protection are some of the security methods that are used. These features make it less important when clients are bad or can't be trusted. It's possible that privacy and security will hurt how useful, clear, fast, and timely the contact is. To make a design safe, privacy and security should not be seen as one-time tech vows, but as constant ways to handle risks.

4.5 Deployment and Governance Dimension

When it comes to sharing and control, the aspect tests whether trustworthy shared intelligence can work with the real edge-cloud limits. For example, edge devices may not have a lot of power, memory, link, or computer power. For example, safe gathering, encryption, differential privacy, finding oddities, and coming up with answers can all raise the cost of computing and communication. Actually, a system that works in theory but not in practice won't be very useful.

Concerns in governance include being able to be audited, who is responsible for what, paperwork, making sure that regulations are followed, rating risk, handling situations that get worse, and keeping people in line. Things like healthcare, banks, public safety, hacking, and infrastructure are especially vulnerable to these fears. With federated learning, technical coordination changes into a responsible social and technical system where decisions can be seen, questioned, and fixed through government.

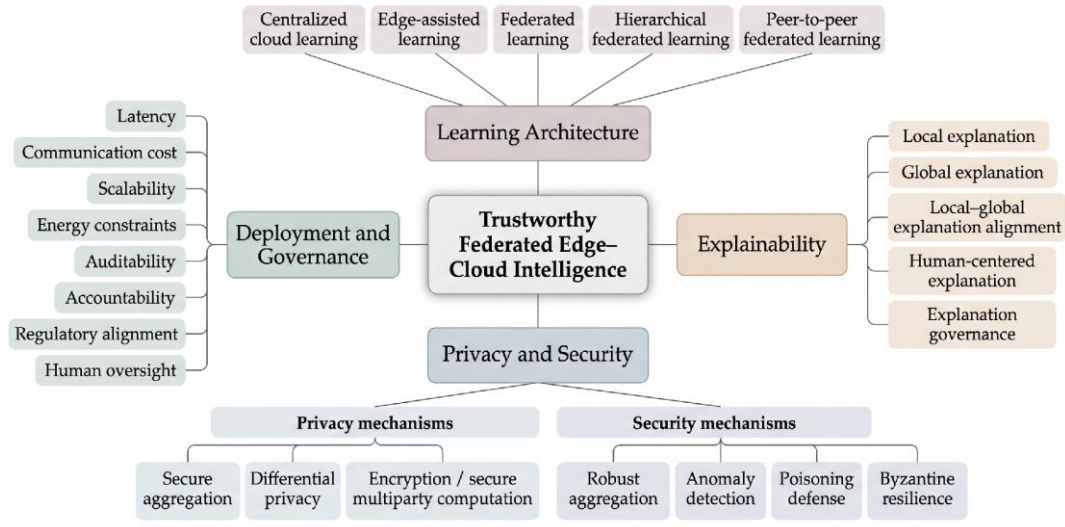


Figure 2. Taxonomy of explainable and trustworthy federated big data intelligence in edge-cloud ecosystems.

Table 3. Taxonomy dimensions and their role in trustworthy federated edge-cloud intelligence.

Taxonomy dimension	Subcategories	Trust function	Example research question
Learning architecture	Federated learning, hierarchical FL, peer-to-peer FL	Determines coordination structure	How should edge clients coordinate without centralizing data?
Explainability	Local, global, local-global alignment, human-centered explanation	Makes distributed decisions understandable	How can explanations be shared without exposing private client information?
Privacy	Differential privacy, secure aggregation, encryption	Reduces direct and indirect data exposure	How can privacy be protected beyond raw-data locality?
Security	Robust aggregation, poisoning defense, anomaly detection	Reduces adversarial influence	How can malicious clients be identified or down-weighted?
Deployment	Latency, bandwidth, scalability, energy use	Supports real-world feasibility	Can trust mechanisms operate under edge constraints?
Governance	Auditability, accountability, oversight, documentation	Supports responsible AI use	Who is responsible for distributed decisions and failures?

5. Proposed Conceptual Framework

5.1 Framework Overview

The suggested approach thinks of reliable shared edge-cloud intelligence as a multi-layered social and technology system instead of a single machine learning method. Throughout the life cycle of a system, trust grows when local data, local analytics, shared cooperation, explainability, privacy protection, security validation, government responsibility, and human control all work together. There are seven linked layers in the system. These are edge data and context, local intelligence, shared cooperation, security and privacy protection, trust governance and

responsibility, and human control. It's not correct to think of these stages as an absolutely straight process. They talk to each other all the time while teaching, deploying, watching, explaining, and going over. Some examples are: a privacy feature could affect the quality of an explanation; a security alert could lead to a review of governance; differences in explanations could mean that clients are leaving; and for human tracking, more audit proof might be needed. Trust is a trait that has many levels and needs to be planned, watched, reviewed, and changed over time.

5.2 Edge Data and Context Layer

Beyond the edge data and context layer, the framework is put together. Industrial systems, healthcare infrastructures, banking platforms, security nodes, users, sensors, devices, and systems that can run on their own are all in it. Near this level, information is usually private, specialized, and only important in the immediate area. The same reading from a monitor could mean different things in different places of work. In the same way, the same clinical signal could mean different things for various patient groups or healthcare facilities. This level of trust depends on the data's quality, where it came from, the agreement's terms, the device's security, how sensitive it is to the situation, and the level of control in the area.

5.3 Local Intelligence Layer

The local intelligence layer trains, thinks, does local analytics, checks, and gives reasons in each client or edge setting. When systems are shared, clients do more than just carry info. They teach the model new things, collect information about their area, and check to see how the model's actions match up with what really happens. This layer is needed to protect the privacy of clients and help people make quick choices in places where that's not possible.

But knowledge from the area makes it harder to believe. It's possible for local models to use data that isn't complete, is out of date, skewed, noisy, or has been changed. Answers from people in the area are good sometimes, but not always. There may be a range of expert skills and control growth levels among clients. In this case, the local intelligence layer needs ways to check the quality of the data, make sure that local explanations are right, stop drift, make sure that changes are safely prepared, and keep track of local assumptions.

5.4 Federated Coordination Layer

Local intelligence and learning in groups are related by the layer of shared management. Parts of this process are choosing clients, communication rounds, syncing, sharing updates, aggregation, and, if needed, multilevel coordination through middle-level edge servers. This layer takes local learning that is spread out and turns it into a global model or shared analysis output. However, the raw data is not kept in one place (McMahan et al., 2017; Kairouz et al., 2021).

Putting things together is more than just math. It's possible that client selection will favor clients with faster connections, bigger datasets, or more powerful computer tools, cutting out smaller or less powerful players. When big trends are brought together, it can be harder to see how minority clients are moving. Customers may be able to see what they're doing or how the system is working during conversation rounds. Since this is the case, a good communication layer should check more than just convergence and speed. It should also check for fairness, privacy, updates that work, inclusion, responsibility, and fairness.

5.5 Explainability and Transparency Layer

Both local and global models can be understood and looked over thanks to the explainability and access layer. In trustworthy shared systems, answers shouldn't come after the model has been taught. You should include them in the shared process so that you can keep an eye on things like how everyone agrees, how explanations change, how accurate explanations are, and how stakeholders can access information over time.

For example, this layer creates alignment signs, stability checks, privacy-risk ratings for explaining things, and reports for stakeholders. For decisions made at the client level, it gives local answers. For actions at the system level,

it gives world reasons. The layer knows that answers need to be kept under control too. They could have private information on them, so access needs to be limited, and there needs to be proof of how they were made and checked.

5.6 Security and Privacy Protection Layer

The element that protects security and privacy works across the whole system. Data about local changes, model changes, contact routes, information, explain results, and aggregation methods are all kept safe. There is still a privacy risk with federated learning because of update leakage, gradient rewriting, membership inference, and explanation leakage (Dwork & Roth, 2014; Bonawitz et al., 2017; Zhu et al., 2019).

Additionally, security controls are very important since compromised clients can send harmful updates, add backdoors, behave in strange ways, or change reasons. Strange thing finding, strong grouping, a client trust score, poisoning defense, and Byzantine resistance can all help lower these risks. There are downsides to these approaches, though. For example, models that are more private might not be as useful, things might take longer if security is tracked more, and more detailed descriptions might make it more likely that information will get out. To avoid ignoring leftover risk, it should be put on paper and watched closely.

5.7 Trust Governance and Accountability Layer

Because of technology problems, the trust control and accountability layer makes sure that the system works in a responsible way. It has things like audit logs, job titles, ways to check on people's work, ways to get help, and rules about who is responsible for what. Within shared systems, various groups are in charge of various parts. People who live nearby, handle data, provide technology, run aggregation services, write models, and make decisions further down the line are all in this group. It is clear how these people can join, what proof they need to give, and how mistakes are looked into within the governance.

Check for client participation, update validation, privacy settings, security alerts, explain outputs, model versions, and human actions. Government papers should explain what the system is built on, what its known boundaries are, what risks are still out there, who is responsible for what, and how to question or fix bad decisions. This method means that the government is built into the technology and can't be separated from it.

5.8 Human Oversight Layer

While automatic coordination methods are used to help people make high-stakes decisions, people are still keeping an eye on them. Someone with oversight must look over the answers with a professional eye, investigate any odd behavior, weigh concerns about fairness, know the differences between local and global systems, and take action when automatic behavior appears to be harmful or wrong. Why something was done, audit logs, risk records, sample paperwork, and clear lines of duty should all be visible to reviewers so that they can do their jobs.

One that is neither symbolic nor just responding is the best kind of human review. Clear levels of development, review methods for big choices, and ways for people with a stake in the outcome to ask questions or challenge the decision should all be built into the system as a whole. Making people more responsible without making shared learning less useful is how having someone watch over things works.

5.9 Cross-Layer Interaction Logic

The idea is based on cross-level association. Edge data stays in the area because it is private, there are legal or policy reasons to do so, or things that make sense. Local models come up with local answers when they learn trends that are specific to each client. Sharing and collecting model updates is how federated coordination turns local learning into information that everyone can use. The ability to explain how local and world models work shows how they work. Privacy and security settings protect information, updates, conversations, and answers. In government, things like risks, decisions, accountability, and who is responsible for what are all put down. Leaders examine facts, respond to signs of danger, and step in when normal teamwork isn't enough. Getting people to trust you takes more than just making one layer better. pooled agreement can still let information get out even if privacy isn't kept safe. You might

get strong but not clear information if you don't talk about security. Explainability that doesn't respect privacy can show private client habits. If you don't have scientific proof, governance can become more about the rules than about what they mean. Human review that isn't based on proof that can be understood and checked may not work. So, the technical, logical, organizational, and human parts must all work together for shared edge-cloud intelligence to be stable.

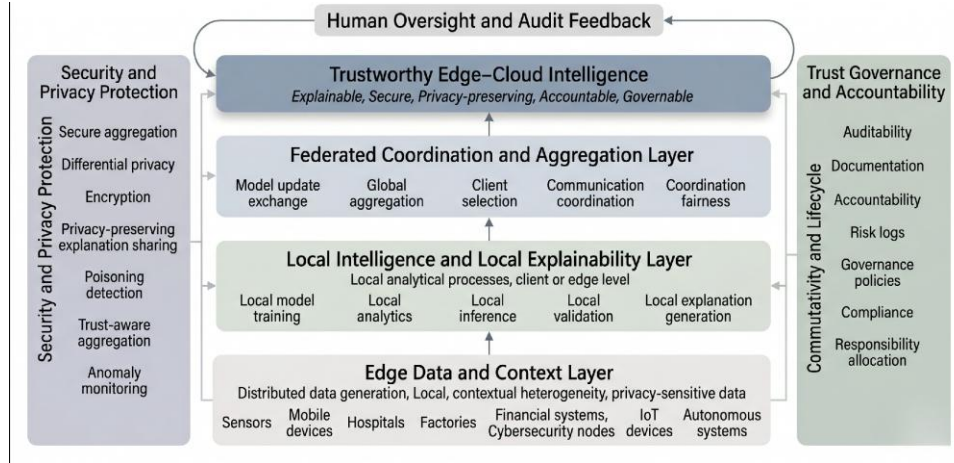


Figure 3. Layered conceptual framework for trustworthy federated edge-cloud intelligence.

Table 4. Layers of the proposed conceptual framework and their functions.

Framework layer	Core function	Trust contribution	Key design question
Edge data and context	Distributed data generation and contextual meaning	Supports data locality and provenance	What data remain local, and why?
Local intelligence	Local training, inference, and explanation	Preserves client autonomy and context sensitivity	How are local models trained, validated, and explained?
Federated coordination	Aggregation of distributed updates	Builds collective intelligence without raw-data pooling	How are updates coordinated fairly and securely?
Explainability and transparency	Local/global explanations and alignment checks	Improves interpretability and contestability	Are explanations faithful, stable, and understandable?
Security and privacy protection	Protection of updates, metadata, communication, and explanations	Reduces leakage and manipulation	What residual risks remain after controls?
Governance and accountability	Audit, documentation, responsibility, compliance	Supports responsible operation	Who is accountable for decisions and failures?
Human oversight	Review, escalation, and intervention	Prevents unchecked automation in high-stakes contexts	When should human reviewers intervene?

6. Analytical Evaluation Dimensions and Governance Criteria

6.1 Technical Evaluation Dimensions

Technical review is still needed because reliable shared intelligence is still needed to make models that are useful and correct. Predictions may be judged by accuracy, F1-score, AUC, loss behavior, and convergence stability in future study that is suitable for the job. But these things don't prove someone is trustworthy on their own. Communication costs, wait times, scalability, energy efficiency, a range of devices, and the ability to work in limited edge situations should all be thought about. There is a chance that a model won't work well in shared edge deployment if it needs a lot of communication rounds, a lot of energy, or delayed inference (Kairouz et al., 2021; Li et al., 2020; Lim et al., 2020; Nguyen et al., 2021).

6.2 Explainability Evaluation Dimensions

The test for explainability should see if it is possible to understand how a common model works at both the local and global levels. It is very important that answers are correct because they should match how the model actually acts, not make sense but aren't true stories. When non-IID data, time shifts, and multiple shared communication rounds happen, it's also important that the rationale stays the same. Going forward, researchers should look at what makes local reasons different from global ones. Leaders examine facts, respond to signs of danger, and step in when normal teamwork isn't enough. People are curious about whether global factors hide risks that are unique to each client. It's important to make sure that the answers are clear for everyone, including subject matter experts, system managers, testers, officials, and users who will be affected. Short or long explanations can slow things down or reveal private information in edge situations. Also, you should think about privacy and explain delays (Adadi & Berrada, 2018; Barredo Arrieta et al., 2020).

6.3 Security and Privacy Evaluation Dimensions

A shared model should be able to be understood at both the local and global levels. This is what the explainability test should look for. Making up stories that make sense but aren't true isn't okay when describing models; they should act like they're described. Additionally, it is important that the reason stays the same when there is non-IID data, time changes, and more than one shared communication round. Scientists should later look into what makes local reasons unique compared to world ones. These people want to know if risks that are unique to each client are hidden by big picture reasons. For this reason, it's important to test how well answers are understood by people who are subject experts, system managers, testers, regulators, and users who will be affected. When answers are too long or too specific, they can slow down processes or reveal private information (Adadi & Berrada, 2018; Barredo Arrieta et al., 2020). This is something that needs to be thought about to protect both privacy and speed.

6.4 Governance and Ethical Evaluation Dimensions

The governance review checks to see if the shared system can be checked, talked about, fought over, and managed in a helpful way. Some people think that we should do more study on law compliance, risk classification, increasing processes, human control, and responsibility systems in the future. Being fair to all clients should also be part of the control rules. This is especially important when non-IID data change the quality or trustworthiness of the model. It's not just technology that's to blame when clients don't do what they're supposed to for different reasons or when different groups are hurt. Society and the government are also to blame. The European Commission High-Level Expert Group on AI (2019) and NIST (2023) both say that a reliable shared system should explain how it learns, why it makes decisions, who is in charge, and how mistakes are found and fixed.

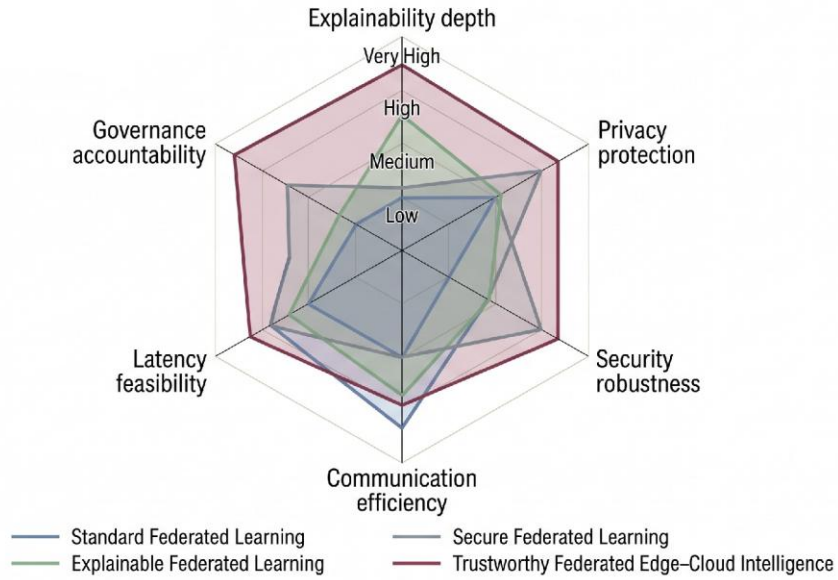


Figure 4. Analytical trade-off map among explainability, privacy protection, security robustness, deployment feasibility, and governance accountability.

Table 5. Evaluation dimensions for future research on trustworthy federated edge-cloud intelligence.

Evaluation category	Example metrics or criteria	Why it matters
Predictive quality	Accuracy, F1-score, AUC, convergence stability	Ensures that the model is useful
Explainability	Fidelity, stability, local-global divergence, interpretability	Ensures that decisions can be understood and reviewed
Privacy	Membership inference risk, leakage score, explanation leakage	Ensures that data are not exposed indirectly
Security	Poisoning resistance, Byzantine tolerance, anomaly detection	Ensures that malicious clients do not dominate learning
Deployment	Latency, bandwidth, energy use, scalability	Ensures real-world edge-cloud feasibility
Governance	Audit logs, documentation, human oversight, accountability allocation	Ensures responsible AI operation
Fairness	Client-level performance disparity, explanation disparity, differential harm	Ensures that minority clients are not ignored

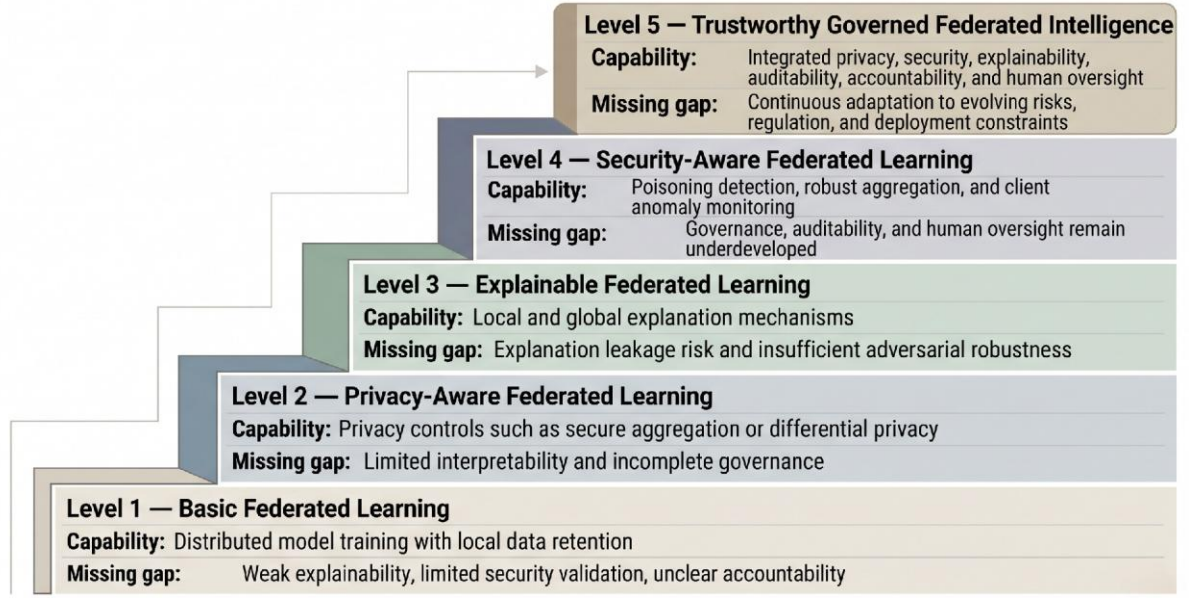


Figure 5. Governance maturity ladder for trustworthy federated edge-cloud intelligence.

7. Research Agenda and Future Directions

Reliable shared edge-cloud intelligence is still a new field of study, as shown by the structure suggestion. It is not yet a set of rules for how things should be done. Even though federated learning has come a long way as a global learning method, there are still some problems that need to be fixed when it comes to privacy, security, being able to explain things, control, fairness, and human tracking. Future research should not ask if models can be taught without keeping raw data in one place. Instead, it should look into how to make distributed intelligence easy to understand, trustworthy, fair, and able to run for a long time.

7.1 Privacy-Preserving Explainability

A big area of study is explainability that keeps private safe in sharing systems. XAI methods that are already in use often assume that they can access centralized data or model behavior. But joint edge-cloud systems make this claim less certain. There may be secret patterns that are only seen by one customer in some accounts. This is especially true when feature attributions show strange behaviors, local flaws, medical signs, financial signals, or operational conditions. Researchers should think about how to gather, compare, and share local replies in the future without giving away private client data. Dwork & Roth (2014) and Bonawitz et al. (2017) say that differently private explanation methods, safe explanation aggregation, encrypted explanation sharing, and privacy-aware explanation checks are all good ideas.

7.2 Causal Explainability in Federated Systems

To describe things in machine learning, most of the time, correlational methods are still used instead of causal ones. Some of them might show which features changed a guess, but not always if those features led to the result or if the model learned to make false links. When customers connect to a shared edge-cloud system, they might see different connections, time conditions, or links depending on where they are. This limitation is key. Our research should look into general causal finding, client-level causal variation, causal explanation agreement, and future-proof causal thinking that protects privacy.

7.3 Standard Benchmarks for Trustworthy Federated Edge-Cloud AI

Slower growth happens over time when marks aren't shared. At the moment, shared learning standards mostly look at data that isn't IID, like how well they can predict success or talk to each other. These tests see how hard it is

to break into something, while XAI tests see how accurate or easy it is to understand something. In this case, all of the usual options are available: non-IID data, privacy risk, harm risk, quality of answers, speed limits, resource limits, fairness, and government review. Some researchers think that shared systems could be used to build trust and make distributed systems better in the future (Kairouz et al., 2021; Li et al., 2020).

7.4 Human-Centered Federated Explainability

An approach to shared explainability that focuses on people has not been developed sufficiently. Different of them need various kinds of information. Changes, quirks, delays, and how accurate updates are may need to be reported to some system officials. Rich people, business owners, doctors, and lawyers are all examples of people who know a lot about a subject and need answers. Some governments may want proof that you are careful, following the rules, and lowering your risks. Damaged users may need clear reasons that let them disagree and trust in a smart way. Researchers like Adadi and Berrada (2018), Barredo Arrieta et al. (2020), Floridi et al. (2018), Guidotti et al. (2018), and Jobin et al. (2019) say that more research is needed to find out if sharing answers helps with knowledge, trust, finding mistakes, quality of control, and staying accountable for choices.

7.5 Trust-Aware and Governance-Aware Aggregation

Fed-based collection shouldn't depend on how many samples are used, how often they are changed, or how well optimization works by itself. If you want to get reliable shared intelligence, you should think about how reliable the clients are, how they handle security, how much privacy is at risk, how stable the reasoning is, how good the data is, how well the rules are followed, and how fair it is for everyone. In the future, researchers should think of ways to combine data that protect real minority-client trends while also lowering the impact of clients who are questionable, unstable, biased, or not well documented. Figure out the difference between bad deviations and good deviations. This is the biggest problem.

7.6 Energy-Aware and Resource-Aware Trustworthy Intelligence

Devices at the edges of networks often have limited power, memory, processing power, and network access. Differential privacy, safe aggregation, encryption, strong validation, anomaly detection, and reason generation are some trust methods that may add extra work. In the future, researchers should work on making trust systems that are small and can work with edge limitations. Some important directions are energy-aware distributed XAI, privacy-preserving updates that are compressed, adaptable communication plans, effective ways to explain things, and trust score that takes into account available resources (Abbas et al., 2018; Lim et al., 2020; Mao et al., 2017; Nguyen et al., 2021).

7.7 Regulatory and Accountability Frameworks

Rules and who is responsible for shared edge-cloud knowledge need to be looked at in more detail in the future. With distributed learning, it's hard to tell who is in charge of what because data owners, local clients, technology providers, managers, developers, and people who make decisions can all change how a system works. Scientists should look into ways to make all of these people responsible, keep track of proof for audits, and make sure that shared systems follow the rules that are unique to each field. Picks that can have a big effect on health, safety, rights, money, or public goods should get extra attention from people.

7.8 Fairness Across Clients and Contexts

Figure out how fair common information is at both the group level and the client level. Although a world model may seem fair on the whole, it may not work well for smaller clients or situations that aren't accurately represented. A similar thing could happen with answers: they might be more reliable for clients who are dominant than for clients who are minority. More research should be done on fairness-aware aggregation, tracking performance at the client level, reason difference, and control methods to stop unfair results in the future. You can't just use numbers to figure out how fair something is; it should also be a part of duty, paperwork, and daily tracking.

8. Conclusion

This article looked at reliable shared edge-cloud intelligence as a problem with many aspects that go beyond training models and predicting outcomes correctly. AI systems work in a wide range of devices, institutions, platforms, and decision-making situations in distributed edge-cloud ecosystems. This makes it hard to make sure that systems are fair, secure, able to be explained, governed, and supervised by humans. The study made an analysis classification that connects control, privacy, security, explainability, learning design, and the ability to apply. Additionally, it suggested a multi-level framework for thinking that would link edge data and context with local intelligence, shared cooperation, being able to explain things and be open about them, security and privacy protection, trust governance and responsibility, and human control. It is clear from this framework that reliability comes from the way that technology protections, interpretive processes, institutional controls, and human review work together.

REFERENCES

1. Abbas, N., Zhang, Y., Taherkordi, A., & Skeie, T. (2018). Mobile edge computing: A survey. *IEEE Internet of Things Journal*, 5(1), 450-465. <https://doi.org/10.1109/JIOT.2017.2750180>
2. Adadi, A., & Berrada, M. (2018). Peeking inside the black-box: A survey on explainable artificial intelligence (XAI). *IEEE Access*, 6, 52138-52160. <https://doi.org/10.1109/ACCESS.2018.2870052>
3. Aledhari, M., Razzak, R., Parizi, R. M., & Saeed, F. (2020). Federated learning: A survey on enabling technologies, protocols, and applications. *IEEE Access*, 8, 140699-140725. <https://doi.org/10.1109/ACCESS.2020.3013541>
4. Bagdasaryan, E., Veit, A., Hua, Y., Estrin, D., & Shmatikov, V. (2020). How to backdoor federated learning. *Proceedings of the 23rd International Conference on Artificial Intelligence and Statistics*, 108, 2938-2948.
5. Barredo Arrieta, A., Diaz-Rodriguez, N., Del Ser, J., Bennetot, A., Tabik, S., Barbado, A., Garcia, S., Gil-Lopez, S., Molina, D., Benjamins, R., Chatila, R., & Herrera, F. (2020). Explainable artificial intelligence (XAI): Concepts, taxonomies, opportunities and challenges toward responsible AI. *Information Fusion*, 58, 82-115. <https://doi.org/10.1016/j.inffus.2019.12.012>
6. Blanchard, P., El Mhamdi, E. M., Guerraoui, R., & Stainer, J. (2017). Machine learning with adversaries: Byzantine tolerant gradient descent. *Advances in Neural Information Processing Systems*, 30.
7. Bonawitz, K., Ivanov, V., Kreuter, B., Marcedone, A., McMahan, H. B., Patel, S., Ramage, D., Segal, A., & Seth, K. (2017). Practical secure aggregation for privacy-preserving machine learning. *Proceedings of the 2017 ACM SIGSAC Conference on Computer and Communications Security*, 1175-1191. <https://doi.org/10.1145/3133956.3133982>
8. Dwork, C., & Roth, A. (2014). *The algorithmic foundations of differential privacy*. Now Publishers. <https://doi.org/10.1561/04000000042>
9. European Commission High-Level Expert Group on Artificial Intelligence. (2019). *Ethics guidelines for trustworthy AI*. European Commission.
10. Floridi, L., Cowls, J., Beltrametti, M., Chatila, R., Chazerand, P., Dignum, V., Luetge, C., Madelin, R., Pagallo, U., Rossi, F., Schafer, B., Valcke, P., & Vayena, E. (2018). AI4People-An ethical framework for a good AI society: Opportunities, risks, principles, and recommendations. *Minds and Machines*, 28, 689-707. <https://doi.org/10.1007/s11023-018-9482-5>
11. Guidotti, R., Monreale, A., Ruggieri, S., Turini, F., Giannotti, F., & Pedreschi, D. (2018). A survey of methods for explaining black box models. *ACM Computing Surveys*, 51(5), Article 93. <https://doi.org/10.1145/3236009>
12. Gunning, D., & Aha, D. (2019). DARPA's explainable artificial intelligence program. *AI Magazine*, 40(2), 44-58. <https://doi.org/10.1609/aimag.v40i2.2850>
13. Jobin, A., Ienca, M., & Vayena, E. (2019). The global landscape of AI ethics guidelines. *Nature Machine Intelligence*, 1, 389-399. <https://doi.org/10.1038/s42256-019-0088-2>
14. Kairouz, P., McMahan, H. B., Avent, B., Bellet, A., Bennis, M., Bhagoji, A. N., Bonawitz, K., Charles, Z., Cormode, G., Cummings, R., D'Oliveira, R. G. L., Eichner, H., El Rouayheb, S., Evans, D., Gardner, J., Garrett, Z., Gascon, A., Ghazi, B., Gibbons, P. B., ... Zhao, S. (2021). Advances and open problems in federated learning. *Foundations and Trends in Machine Learning*, 14(1-2), 1-210. <https://doi.org/10.1561/22000000083>
15. Li, T., Sahu, A. K., Talwalkar, A., & Smith, V. (2020). Federated learning: Challenges, methods, and future directions. *IEEE Signal Processing Magazine*, 37(3), 50-60. <https://doi.org/10.1109/MSP.2020.2975749>
16. Li, T., Sahu, A. K., Zaheer, M., Sanjabi, M., Talwalkar, A., & Smith, V. (2020). Federated optimization in heterogeneous networks. *Proceedings of Machine Learning and Systems*, 2, 429-450.
17. Lim, W. Y. B., Luong, N. C., Hoang, D. T., Jiao, Y., Liang, Y.-C., Yang, Q., Niyato, D., & Miao, C. (2020). Federated learning in mobile edge networks: A comprehensive survey. *IEEE Communications Surveys & Tutorials*, 22(3), 2031-2063. <https://doi.org/10.1109/COMST.2020.2986024>
18. Mao, Y., You, C., Zhang, J., Huang, K., & Letaief, K. B. (2017). A survey on mobile edge computing: The communication perspective. *IEEE Communications Surveys & Tutorials*, 19(4), 2322-2358. <https://doi.org/10.1109/COMST.2017.2745201>

19. McMahan, H. B., Moore, E., Ramage, D., Hampson, S., & Aguera y Arcas, B. (2017). Communication-efficient learning of deep networks from decentralized data. *Proceedings of the 20th International Conference on Artificial Intelligence and Statistics*, 54, 1273-1282.
20. Mittelstadt, B. (2019). Principles alone cannot guarantee ethical AI. *Nature Machine Intelligence*, 1, 501-507. <https://doi.org/10.1038/s42256-019-0114-4>
21. Mothukuri, V., Parizi, R. M., Pouriyeh, S., Huang, Y., Dehghantanha, A., & Srivastava, G. (2021). A survey on security and privacy of federated learning. *Future Generation Computer Systems*, 115, 619-640. <https://doi.org/10.1016/j.future.2020.10.007>
22. Nasr, M., Shokri, R., & Houmansadr, A. (2019). Comprehensive privacy analysis of deep learning: Passive and active white-box inference attacks against centralized and federated learning. *2019 IEEE Symposium on Security and Privacy*, 739-753. <https://doi.org/10.1109/SP.2019.00065>
23. National Institute of Standards and Technology. (2023). Artificial intelligence risk management framework (AI RMF 1.0). U.S. Department of Commerce. <https://doi.org/10.6028/NIST.AI.100-1>
24. Nguyen, D. C., Ding, M., Pathirana, P. N., Seneviratne, A., Li, J., & Poor, H. V. (2021). Federated learning for Internet of Things: A comprehensive survey. *IEEE Communications Surveys & Tutorials*, 23(3), 1622-1658. <https://doi.org/10.1109/COMST.2021.3075439>
25. Ribeiro, M. T., Singh, S., & Guestrin, C. (2016). "Why should I trust you?" Explaining the predictions of any classifier. *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, 1135-1144. <https://doi.org/10.1145/2939672.2939778>
26. Rieke, N., Hancox, J., Li, W., Milletari, F., Roth, H. R., Albarqouni, S., Bakas, S., Galtier, M. N., Landman, B. A., Maier-Hein, K., Ourselin, S., Sheller, M., Summers, R. M., Trask, A., Xu, D., Baust, M., & Cardoso, M. J. (2020). The future of digital health with federated learning. *npj Digital Medicine*, 3, Article 119. <https://doi.org/10.1038/s41746-020-00323-1>
27. Roman, R., Lopez, J., & Mambo, M. (2018). Mobile edge computing, fog et al.: A survey and analysis of security threats and challenges. *Future Generation Computer Systems*, 78, 680-698. <https://doi.org/10.1016/j.future.2016.11.009>
28. Satyanarayanan, M. (2017). The emergence of edge computing. *Computer*, 50(1), 30-39. <https://doi.org/10.1109/MC.2017.9>
29. Shi, W., Cao, J., Zhang, Q., Li, Y., & Xu, L. (2016). Edge computing: Vision and challenges. *IEEE Internet of Things Journal*, 3(5), 637-646. <https://doi.org/10.1109/JIOT.2016.2579198>
30. Xu, J., Glicksberg, B. S., Su, C., Walker, P., Bian, J., & Wang, F. (2021). Federated learning for healthcare informatics. *Journal of Healthcare Informatics Research*, 5, 1-19. <https://doi.org/10.1007/s41666-020-00082-4>
31. Yang, Q., Liu, Y., Chen, T., & Tong, Y. (2019). Federated machine learning: Concept and applications. *ACM Transactions on Intelligent Systems and Technology*, 10(2), Article 12. <https://doi.org/10.1145/3298981>
32. Zhang, C., Xie, Y., Bai, H., Yu, B., Li, W., & Gao, Y. (2021). A survey on federated learning. *Knowledge-Based Systems*, 216, Article 106775. <https://doi.org/10.1016/j.knosys.2021.106775>
33. Zhu, L., Liu, Z., & Han, S. (2019). Deep leakage from gradients. *Advances in Neural Information Processing Systems*, 32.