

# Viabilidade do aprendizado federado para desenvolvimento de soluções de IA em diagnóstico laboratorial sob perspectiva de privacidade: uma revisão sistemática

Rodrigo Cunha Dias<sup>1</sup>, Márcia Ito<sup>12</sup>

<sup>1</sup> Programa de Educação Continuada da Escola Politécnica de USP (PECE)  
Universidade de São Paulo (USP)  
São Paulo, SP – Brasil

<sup>2</sup> Programa de Mestrado Profissional em Sistemas Produtivos  
Centro Estadual de Educação Tecnológica Paula Souza (CEETEPS)  
São Paulo, SP – Brasil

{qualit.rcd@uol.com.br; marcia.ito01@cps.sp.gov.br}

**Abstract.** *The development of artificial intelligence solutions in laboratory diagnostics faces a structural dilemma: the requirement for large data volumes for training conflicts with strict health data protection regulations (GDPR, LGPD, HIPAA). This systematic review investigated federated learning as a methodological alternative that enables AI model training without centralizing sensitive data. Analyzing 23 articles (2021-2025), we found that 56.5% of studies demonstrated technical feasibility, although only 26.1% explicitly addressed regulatory compliance. Results reveal a technically promising field, but with a critical gap between technical maturity and regulatory adequacy.*

**Keywords.** *Federated Learning; Laboratory Diagnosis; Data Privacy; Artificial Intelligence in Healthcare; LGPD (Brazilian General Data Protection Law); Systematic Review.*

**Resumo.** *O desenvolvimento de soluções de inteligência artificial em diagnóstico laboratorial enfrenta um dilema estrutural: a exigência de grandes volumes de dados para treinamento contrapõe-se às rigorosas legislações de proteção de dados em saúde (LGPD, GDPR, HIPAA). Esta revisão sistemática investigou o aprendizado federado como alternativa metodológica que possibilita o treinamento de modelos de IA sem centralização de dados sensíveis. Analisando 23 artigos (2021-2025), identificou-se que 56,5% dos estudos demonstraram viabilidade técnica, embora apenas 26,1% tenham abordado conformidade regulatória. Os resultados revelam um campo tecnicamente promissor, mas com gap crítico entre maturidade técnica e adequação regulatória.*

**Palavras-Chave.** *Aprendizado Federado; Diagnóstico Laboratorial; Privacidade de Dados; Inteligência Artificial em Saúde; LGPD (Lei Geral de Proteção de Dados); Revisão Sistemática.*

## 1. Introdução

A medicina contemporânea vivencia uma busca contínua por melhores desfechos clínicos, impulsionando o desenvolvimento de novas formas de diagnóstico e prognóstico. Nesse cenário, a tecnologia atua como um catalisador fundamental, com destaque para a Inteligência Artificial (IA), que recebe atenção especial devido à sua capacidade de processar volumes de dados frequentemente incompreensíveis à mente humana.

Embora a aplicação de algoritmos avançados já esteja mais desenvolvida em algumas áreas da medicina diagnóstica, o segmento de análises laboratoriais, detentor de vastos repositórios de dados, começa agora a despertar para o potencial transformador dessas ferramentas. Contudo, a transposição do potencial teórico da IA para a prática laboratorial enfrenta obstáculos que vão além da simples adoção tecnológica. Questões como a governança da qualidade dos dados e a necessidade de novas competências profissionais, que unam o conhecimento clínico laboratorial ao pensamento computacional, impõem-se como requisitos básicos para o sucesso desses projetos. Mas é na tensão entre a necessidade de dados para o treinamento dos algoritmos e a obrigatoriedade da proteção das informações dos pacientes que reside o desafio mais crítico e estrutural para o setor.

O modelo tradicional de desenvolvimento de IA depende, invariavelmente, da centralização de dados. Ao transferir informações de suas bases de origem para repositórios centrais de treinamento, cria-se dessa forma um cenário de vulnerabilidade que eleva significativamente os riscos de exposição de dados sensíveis. Em contrapartida, os marcos regulatórios de proteção de dados, como a Lei Geral de Proteção de Dados (LGPD), estabelecem barreiras de segurança indispensáveis e necessárias que, na prática, tornam o acesso a essas informações extremamente restrito e burocrático. Forma-se, assim, um impasse: a movimentação dos dados é arriscada demais sob a ótica da privacidade, enquanto o bloqueio do acesso inviabiliza a inovação tecnológica necessária dado a velocidade das transformações tecnológicas.

Diante da necessidade de superar essa dicotomia, surgiu o Aprendizado Federado como uma mudança de paradigma metodológico. Essa abordagem subverte a lógica convencional ao evitar a transferência dos dados e permitir que o treinamento ocorra localmente. Dessa forma, teoricamente é possível desenvolver modelos robustos de inteligência artificial de maneira colaborativa, sem que os dados sensíveis dos pacientes jamais precisem sair do ambiente seguro do laboratório, alinhando a inovação tecnológica às mais estritas exigências de conformidade e ética.

Considerando esse panorama de oportunidades emergentes e desafios regulatórios, este trabalho tem como objetivo identificar e analisar evidências sobre a aplicação do aprendizado federado como alternativa viável para o desenvolvimento de soluções de IA em diagnóstico laboratorial, avaliando benefícios, limitações e desafios quanto à preservação da privacidade de dados sensíveis e à conformidade regulatória em saúde.

### 1.2 Uso de Inteligência Artificial em Laboratórios Médicos

A aplicação de IA na medicina diagnóstica teve grande desenvolvimento em exames por imagem e chegou ao segmento de análises laboratoriais onde há vasta quantidade de dados estruturados [Kodali et al. 2024]. Devido às suas peculiaridades, o laboratório

médico reúne características na estruturação de seus dados que o torna forte candidato a obter grande sucesso em aplicação de soluções de inteligência artificial [Xiao 2023]. No entanto, desafios existem e merecem atenção especial.

Escolher a técnica adequada para cada necessidade do contexto laboratorial não é tarefa simples. Trata-se de uma questão complexa que envolve o desconhecimento dos fundamentos teóricos dos modelos por muitos que atuam no segmento saúde e a necessidade de rigoroso processo de desenvolvimento, documentação e validação antes de sua aplicação. Segundo Cadamuro et al. (2024), os profissionais dos laboratórios europeus apontaram a necessidade de novas competências, aprofundamento de domínios matemáticos e aspectos éticos-regulatórios como as principais barreiras a serem transpostas para a adoção de IA nos laboratórios médicos. Corroboram com essa percepção Adler et al. (2023), que salientam existir deficiências na análise de dados de forma geral, sugerindo a adoção de temas relacionados ao pensamento computacional nas futuras formações acadêmicas dos profissionais de medicina laboratorial. Para Prabh et al. (2024), a falta de conhecimento técnico-científico ao escolher o algoritmo adequado aplicável à cada demanda laboratorial é um aspecto a ser considerado.

Outra questão crítica é o nível de qualidade dos dados que pode ser comprometida por fatores operacionais, falta de padrão e falha humana. Segundo Iturry et al (2021), ao realizarem revisão de literatura concluíram que os estudos sobre a qualidade dos dados em saúde apontavam como problemas mais relevantes aqueles relacionados às dimensões completude, correção, consistência e precisão. A baixa qualidade pode introduzir viés ao modelo ou inviabilizar o uso dos dados para desenvolvimento de soluções de inteligência artificial nessas organizações. A generalização dos dados, considerando dados de diferentes serviços laboratoriais, é tido como desejável para minimizar viés algorítmico relacionado às características institucionais.

### 1.3 Contexto Brasileiro da Inteligência Artificial em Laboratórios Médicos

O uso de soluções de inteligência artificial em laboratórios médicos é muito insipiente em todo o mundo. No contexto brasileiro, os laboratórios médicos acompanham a evolução da adoção de soluções de IA nos processos laboratoriais com muita atenção e interesse em razão de seu potencial transformador na área diagnóstica. Por isso, a adoção da inteligência artificial na saúde vem progressivamente aumentando e novas aplicações e casos de uso surgem a cada dia.

O crescente interesse dos laboratórios médicos sobre o tema indica que essa fase exponencial de uso de IA em saúde continuará em ascensão tornando essa tecnologia gradativamente, cada vez mais, presente no cotidiano dos laboratórios [Lippi and Plebani 2025].

No entanto, existe carência de estudos quanto ao cenário de adoção de IA nos laboratórios médicos brasileiros.

Um relatório de mercado gerado pela Associação Brasileira de Medicina Diagnóstica (ABRAMED), circunscrito a seus associados que correspondem a 65% dos exames realizados na saúde suplementar brasileira (segundo o relatório), indicou que mensalmente de 10 a 140 mil pacientes são impactados por IA nos serviços recebidos, direta ou indiretamente em seus resultados de exames [ABRAMED 2024].

A Federação Internacional de Química Clínica e Medicina Laboratorial (IFCC) publicou recomendações para aplicação de aprendizado de máquina em laboratórios médicos. Foi dada especial atenção para a necessidade de haver raciocínio lógico guiado por conhecimentos sobre a natureza dos dados, preocupações éticas e planejamento meticuloso baseado em conhecimento sobre quais tipos de problemas merecem, ou não, serem tratados por meio de aprendizado de máquina [Master et al. 2023].

#### 1.4 Proteção de Dados e Regulação

Para treinar modelos de IA, dados são necessários. No contexto de saúde, dados sensíveis. A quantidade de dados gerados diariamente no mundo cresce de forma exponencial e o segmento saúde tem acompanhado o mesmo fenômeno *big data*. Logo, a questão não é escassez de dados, mas sim o acesso a eles. No caso, os dados são protegidos por legislação específica para preservar a privacidade das pessoas naturais, como a Lei Geral de Proteção de Dados (LGPD) brasileira, o Regulamento Geral de Proteção de dados (GDPR) europeu e a *Health Insurance Portability and Accountability Act* (HIPAA) dos EUA. Assim, ao mesmo tempo que muitas oportunidades se abrem para que avanços sejam alcançados através do uso das soluções de IA em saúde, faz-se necessário ter cuidados especiais para que não haja exposição de dados sensíveis. Essa preocupação deve ser elemento central, balizando o uso ético e seguro dos dados em saúde em medicina laboratorial.

Considerando a forma tradicional de desenvolvimento de modelos de IA, devido ao fato desta ser baseada em transferência e centralização de dados, ficou evidente a crescente ameaça à violação dos dados através desse método. Sendo assim, investir recursos para evitar quebra de privacidade de dados em saúde vem sendo crescentemente utilizado, sendo comum a atuação sinérgica das áreas que compõem as organizações de saúde para evitar comprometimento dos seus dados, buscando evitar sinistros de privacidade e mantendo o foco no direito dos seus pacientes.

Como já mencionado, privacidade de dados é uma questão fortemente regulamentada com pesadas sanções às organizações que se veem com dados vazados, além do aspecto reputacional que certamente afetará suas relações com as partes interessadas. Por outro lado, as restrições cada vez mais exigentes de acesso aos dados de saúde, em alguns casos chegando a inviabilizar o desenvolvimento das soluções de IA pode impedir a evolução à novas soluções em saúde, essa questão foi considerada por Cadamuro et al. (2024). Os autores complementam argumentando que apesar de haver plena compreensão sobre a necessidade da salvaguarda dos dados sensíveis em saúde dos pacientes, as travas legais de segurança para a proteção desses dados levam à demasiada demora em desenvolvimento que, frente aos rápidos avanços tecnológicos, pode trazer prejuízos à melhoria em saúde circunscrito ao perímetro regional da abrangência legal.

A proteção de dados sensíveis em saúde começa a receber atenção no contexto brasileiro por meio da Política Nacional de Informação e Informática em Saúde (PNIIS) sendo sua primeira versão em 2004. No entanto, o principal marco legal referente à privacidade de dados no cenário brasileiro foi a Lei 13.709 de 2018, Lei Geral de Proteção de Dados Pessoais (LGPD), que foi inspirada na General Data Protection Regulation (GDPR) da União Europeia. A fiscalização do cumprimento da LGPD é de competência da Agência Nacional de Proteção de Dados (ANPD).

Considerando a governança dos dados em saúde, esse papel ficou a cargo da Rede Nacional de Dados em Saúde (RNDS). Sendo assim, sob seu escopo de atuação ficou a definição de aspectos técnicos necessários para gerir os dados de saúde no ambiente digital, definindo quais padrões para interoperabilidade (protocolos HL7 FHIR), qual tecnologia adotar para proteção (criptografia assimétrica), como controlar e registrar acesso aos dados pelo titular e pelos profissionais de saúde e suas organizações. No entanto, a natureza da RNDS é a proteção de dados do cidadão [Fantonelli et al. 2020]. Logo, as travas naturalmente existentes para acessar os dados se tornariam mais rígidas podendo ser inviáveis para realizar treinamentos de modelos de inteligência artificial. Cabe salientar que a escolha do padrão HL7 FHIR carrega consigo grande ganho e vantagens por permitir integração com protocolos difundidos internacionalmente. Por outro lado, apresenta desafios elevados em sua adoção prática por envolver maior complexidade, levando à demanda por suportes técnicos significativos e, por isso, o efeito sobre a interoperabilidade de dados na saúde poderá ser ampliado, resultando em maior desafio à sua implementação real e prática [Fantonelli et al. 2020].

Nesse cenário, dos laboratórios médicos brasileiros, surge a necessidade de explorar estratégias complementares como o aprendizado federado para que os laboratórios que queiram desenvolver soluções e/ou modelos de IA o façam de forma segura em relação à privacidade de dados de seus pacientes, respeitando os requisitos legais aplicáveis.

Nesse *trade-off*, onde tem-se de um dos lados a necessidade inquestionável da privacidade de dados sensíveis em saúde e, do outro, o anseio por contribuições de alto valor à saúde oriundos das soluções baseadas em inteligência artificial, é que surge o Aprendizado Federado (AF). Em comparação ao método tradicional de treinamento de IA com centralização de dados, o AF apresenta vantagens distintas de privacidade [McMahan et al. 2017].

## 1.5 O aprendizado federado

O aprendizado federado (AF) foi introduzido por [McMahan et al. (2017)] por meio do algoritmo *Federated Averaging* (FedAvg). Essa proposta representa uma mudança metodológica profunda em relação ao modelo tradicional de treinamento de sistemas de IA. A principal diferença é a ausência de transferência de dados entre as instituições participantes, de maneira que cada conjunto de dados permanece armazenado exclusivamente em seu local de origem. Essa característica do AF reduz significativamente os riscos associados à privacidade de dados sensíveis, pois apenas as atualizações produzidas durante o treinamento local, não os dados em si, são enviadas ao servidor responsável pela agregação global. Dessa forma, instituições distintas podem colaborar na construção de modelos mais robustos sem renunciar à confiabilidade e controle sobre seus dados sensíveis.

No AF, os diferentes locais onde dados são gerados e armazenados são denominados clientes. Esses clientes participam de uma rede federada coordenada por um servidor central. Durante o treinamento, o servidor envia o modelo global para os clientes, os quais realizam o ajuste local dos parâmetros utilizando unicamente seus próprios dados. Ao final dessa etapa, cada cliente retorna uma atualização, que será incorporada ao modelo geral por meio de uma regra de agregação ponderada. Esse processo é iterativo, se repetindo até que os ganhos obtidos entre uma rodada e outra se

estabilizem, o que está diretamente relacionado ao problema de otimização buscando minimizar a função objetivo global.

Segundo a formulação teórica apresentada no artigo seminal de [McMahan et al \(2017\)](#), o objetivo do AF consiste em encontrar o menor valor possível da função objetivo global representada na Figura 1 pela Fórmula (1):

$$f(w) \stackrel{\text{def}}{=} \frac{1}{n} \sum_{i=1}^n f_i(w) \quad (1)$$

Nesse contexto,  $w$  representa o vetor de parâmetros do modelo, pertencente ao espaço vetorial dos números reais. Essa expressão define o problema de otimização a ser resolvido: minimizar a média das perdas individuais avaliadas em todos os dados distribuídos entre os clientes. O vetor  $w$  contém todos os parâmetros ajustáveis do modelo e sua atualização iterativa busca reduzir  $f(w)$ , ainda que os dados estejam descentralizados e distribuídos de forma heterogênea.

Durante o treinamento local, cada cliente estima sua própria função de perda agregada, descrita pela Fórmula (2):

$$F_k(w) = \frac{1}{n_k} \sum_{i \in P_k} f_i(w) \quad (2)$$

Na qual  $F_k(w)$  corresponde à perda média calculada apenas sobre os exemplos pertencentes ao cliente  $k$ , e  $n_k$  representa a quantidade de dados armazenados localmente. Essa formulação reflete explicitamente o fato de que cada cliente possui um conjunto distinto de exemplos e que os dados não são independentes nem identicamente distribuídos (não-IID). No caso de laboratórios médicos, essa característica é particularmente evidente, já que cada instituição atende populações distintas e coleta dados sob condições próprias, resultando em conjuntos desbalanceados e heterogêneos.

Após receber as perdas locais  $F_k(w)$ , o servidor realiza a etapa de agregação global, definida pela Fórmula (3), também ilustrada na Figura 1:

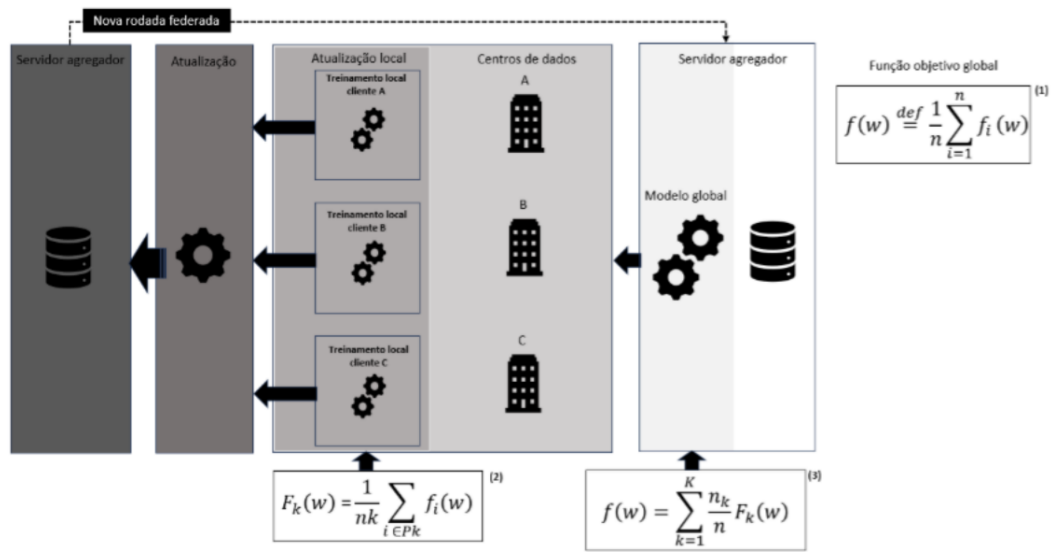
$$f(w) = \sum_{k=1}^K \frac{n_k}{n} F_k(w) \quad (3)$$

Essa equação estabelece que cada atualização local deve contribuir para o modelo global de forma ponderada, proporcional ao número de dados que o cliente possui. Assim, clientes com bases de dados maiores influenciam mais intensamente a atualização global, garantindo representatividade proporcional e evitando distorções que poderiam ser originadas caso todos tivessem peso igual. É justamente esse mecanismo que permite que



o modelo global se beneficie da diversidade dos dados distribuídos, sem que haja transferência de informações sensíveis entre instituições.

Um ciclo completo abrangendo envio do modelo global, treinamento local, agregação e atualização global é denominado rodada federada conforme descrito por [Sheller et al. \(2020\)](#). A Figura 1 ilustra essa dinâmica, evidenciando a circulação do modelo entre o servidor e os clientes, bem como a relação entre as fases do processo e as formulações matemáticas, representando respectivamente: a função objetivo global, a perda local calculada pelos clientes e a regra de agregação global responsável pela atualização do modelo.



**Figura 1. Rodada federada**

Dessa forma, a integração entre o diagrama apresentado na Figura 1 e as questões fundamentais do AF evidenciam a lógica distribuída, iterativa e colaborativa do método. Esse modelo oferece um mecanismo robusto para a construção de sistemas de IA em ambientes onde a centralização de dados é impraticável, como em laboratórios médicos. A heterogeneidade natural dessas instituições, volume de dados, padrão populacional, método de coleta, não viola o paradigma federado; ao contrário, constitui exatamente o tipo de cenário para o qual ele foi concebido. Sendo assim, o AF reúne as características para torná-lo viável e alinhado às premissas de privacidade de dados em saúde, aplicável ao desenvolvimento de modelos de IA no segmento diagnóstico dos laboratórios médicos.

Apesar do AF conferir importante avanço referente à proteção da privacidade dos dados por mantê-los em seus locais de origem, transmitindo somente a atualização do modelo; ele não está salvo de ameaças e ataques. Vários estudos recentes demonstram que existem ameaças ao AF que são capazes de comprometer a proteção à privacidade dos dados como o ataque de inferência na atualização do gradiente, reconstrução de dados, envenenamento de modelo e ataques bizantinos [\[Guembe et al. 2024\]](#), [\[Jimenez-Gutierrez et al. 2025\]](#), [\[Feng et al. 2025\]](#), [\[Zhao et al. 2025\]](#).

## 2 Método

Realizada uma revisão sistemática de literatura (RSL) segundo [Kitchenham et al. \(2007\)](#) e baseada no protocolo PRISMA [\[Page et al. 2021\]](#). A RSL é um estudo do tipo secundário elaborado de forma metódica seguindo rigorosa sistemática em sua elaboração e execução, permitindo sua reprodução. Trabalha com evidências científicas disponíveis referente ao tema de pesquisa, consultando-as de forma estruturada à luz da questão de pesquisa selecionada.

Atualmente existem muitas soluções tecnológicas disponíveis para auxiliar nesse tipo de estudo, foi escolhida a ferramenta online Parsifal (versão 2.2). Sua aplicação no estudo se deu nas fases de planejamento, condução e elaboração do relatório da revisão.

### 2.1 Planejamento

A fase planejamento foi desenvolvida segundo [Kitchenham et al. \(2007\)](#), sendo elaborada por meio do uso do sistema online Parsifal (versão 2.2). O objetivo da revisão sistemática da literatura foi definido como: identificar e analisar evidências científicas sobre a aplicação do aprendizado federado como alternativa viável para o desenvolvimento de soluções de IA em diagnóstico laboratorial, avaliando benefícios, limitações e desafios quanto à preservação da privacidade de dados sensíveis e à conformidade regulatória em saúde.

### 2.2 Estratégia PICOC e Perguntas de Pesquisa

Partindo do objetivo proposto, foi feito uso da estratégia PICOC levando ao desdobramento do objetivo de maneira a compreender o relacionamento do fenômeno estudado e as dimensões população, intervenção, comparação, desfecho e contexto. Dessa forma houve a condução às perguntas de pesquisa. As questões de pesquisa (QP) permitem aplicar seleção sistemática dos artigos nas fontes de dados, guardando a estreita relação com o objetivo definido para o trabalho. Dessa forma, foram elaboradas as três Questões de Pesquisa (QP):

QP1: O aprendizado federado é uma alternativa para o desenvolvimento de soluções de IA seguras no diagnóstico laboratorial?

QP2: Quais são os principais desafios técnicos, operacionais e regulatórios para sua adoção em ambientes laboratoriais?

QP3: Como o aprendizado federado é comparado aos modelos centralizados de IA sob a ótica da privacidade e segurança de dados em saúde?

As três questões de pesquisa foram desdobradas em seis Questões para Avaliação da Qualidade (QAQ) dos artigos, guardando alinhamento com o objetivo da RSL. Essas seis QAQ foram aplicadas na fase condução com o propósito de classificar os artigos como elegíveis ao estudo:

QAQ 1.1: O estudo apresenta evidências empíricas (experimentos, métricas ou estudos de caso) que demonstram a viabilidade ou eficácia do aprendizado federado em diagnóstico laboratorial ou contexto equivalente de saúde?

QAQ 1.2: A pesquisa demonstra resultados mensuráveis de desempenho, privacidade ou segurança que evidenciam benefícios do aprendizado federado sobre abordagens tradicionais?



QAQ 2.1: O artigo descreve desafios técnicos e operacionais (ex.: heterogeneidade de dados, comunicação, recursos computacionais) que afetam a implementação do aprendizado federado em ambientes laboratoriais?

QAQ 2.2: O estudo considera barreiras regulatórias, éticas ou de conformidade (ex.: LGPD, GDPR, confidencialidade médica) na aplicação do aprendizado federado à saúde?

QAQ 3.1: O artigo realiza comparações explícitas com modelos centralizados de IA avaliando aspectos de desempenho, privacidade ou eficiência?

QAQ 3.2: As conclusões do estudo abordam implicações para a segurança e privacidade de dados em saúde, discutindo vantagens e limitações do aprendizado federado frente a abordagens centralizadas?

Para cada artigo que tenha atendido aos critérios de elegibilidade, foi adotado um critério de score para classificar as respostas às QAQ. Para responder as QAQ foram considerados o título do artigo, palavras chaves declaradas e o resumo do artigo descrito por seus autores.

### 2.3 Estratégia de Busca e Palavras-Chaves

A estratégia PICOC também ajuda no raciocínio e escolha das palavras chaves para a estratégia de busca, pois permite contextualizar o objetivo, aprofunda a abordagem e passa por diferentes camadas de análise que reposiciona o problema conforme as diferentes dimensões. Como resultado, aguça-se a compreensão sobre o objetivo proposto e surgem palavras capazes de sintetizar a problemática proposta. As palavras chaves foram definidas como: (a) aprendizado federado em diagnóstico laboratorial; (b) IA em saúde; (c) aprendizado federado; (d) diagnóstico laboratorial; (e) privacidade de dados em saúde.

Partindo das palavras chaves e seus sinônimos, considerando ainda a maneira como a literatura acadêmica habitualmente se refere aos temas, foi elaborada a estratégia de busca nas bases de dados, a saber:

("federated learning" OR "distributed learning") AND ("clinical laboratory" OR "laboratory medicine" OR "laboratory diagnostics") AND ("data privacy" OR "data protection" OR "privacy-preserving") AND ("healthcare").

As buscas foram realizadas nas bases de dados: IEEE Digital Library, Scopus e Springer Link.

### 2.4 Critérios de Inclusão e Exclusão

Como critério de inclusão foram definidos: artigos completos e revisados por pares; documentos regulatórios; estudos no idioma inglês e português; estudos primários ou experimentais que abordem aprendizado federado aplicado à saúde; estudos que descrevam desafios, implementações, resultados ou impactos do aprendizado federado na privacidade e segurança de dados em saúde; publicações no período de 2021 a 2025; textos completos, revisados por pares, disponíveis em inglês ou português.

A delimitação do idioma incluindo o inglês, se justificada pela predominância dessa língua na literatura científica indexada sobre aprendizado federado e diagnóstico laboratorial, no entanto, pode introduzir um viés de idioma (*language bias*), na medida em que estudos relevantes publicados em outras línguas (ex.: chinês, alemão, espanhol) podem ter sido sistematicamente excluídos.

Como critério de exclusão foram definidos: aplicações de IA em saúde que não tratem de privacidade ou proteção de dados; estudos fora do contexto da saúde humana; idiomas diferentes do inglês e português; livros; publicações duplicadas, incompletas, revisões ou sem acesso integral; trabalhos sem relação direta com aprendizado federado.

## 2.5 Fase Condução

Nessa fase iniciaram as buscas pelos artigos e sua seleção através da aplicação dos critérios definidos na fase planejamento sendo usado o sistema Parsifal. A fase condução possui as etapas pesquisa, importação, seleção, avaliação da qualidade, extração de dados e análise dos dados.

A pesquisa foi feita usando a estratégia de busca definida na fase planejamento aplicada nas fontes de dados em seus respectivos websites: IEEE Digital Library, Scopus e Springer Link. Os termos de busca e os operadores booleanos foram usados sem otimização para as diferentes fontes, ou seja, foram usados da mesma forma para todas as bases. A pesquisa foi efetuada no mês de outubro de 2025 e foram considerados os seguintes critérios para todas as fontes: período de busca por ano com início em 2021 até 2025 e exclusão de artigos de revisão.

A importação foi realizada gerando arquivos com todos os artigos resultantes da busca em cada fonte, os arquivos foram importados para o sistema Parsifal.

A seleção preliminar ocorreu aplicando os critérios de inclusão e exclusão definidos no planejamento da pesquisa o que gerou o grupo de artigos classificados como Artigos Grupo 1 (AG1). Em seguida, foram aplicadas as QAQ apenas nos títulos, resumos e palavras chaves de cada artigo AG1 adotando um critério de score de pontuação explicado a seguir. Como resultado foi gerado novo subgrupo de Artigos Grupo 2 (AG2). Nesse subgrupo AG2 foram aplicadas as QED em cada artigo considerando todo o seu conteúdo para implementar a estratégia de síntese analítica multinível detalhada a frente.

A análise dos dados foi realizada com base na matriz booleana gerada a partir da estratégia de síntese analítica multinível.

Todas as etapas descritas acima estão detalhadas nas próximas sessões.

## 2.6 Questão de Análise da Qualidade (QAQ) – critério de aplicação

Os artigos do AG1 foram avaliados em relação à sua qualidade considerando o objetivo da RSL. Isso se deu por meio da aplicação das QAQ no AG1 o que resultou na classificação dos Artigos Grupo 2 (AG2).

A sistemática considerou os títulos dos artigos, seus resumos e palavras-chaves para responder as seis QAQ (QAQs 1.1, 1.2, 2.1, 2.2, 3.1 e 3.2) em cada artigo do AG1 buscando encontrar respostas para cada questão.

Foi aplicado um critério de pontuação obedecendo a seguinte lógica para a capacidade de cada artigo responder a cada uma das seis QAQ:

- a) Atende totalmente (peso 2.0): o artigo atende todos os aspectos da questão;
- b) Atende parcialmente (peso 1.0): o artigo atende um ou mais aspectos da questão, não todos;
- c) Não atende (peso 0.0): o artigo não atende os aspectos da questão.

O score máximo foi definido em 12 pontos e o *cutoff score* em 8.0 pontos.

Sendo assim, os Artigos Grupo 2 (AG2) formaram a seleção de artigos que pertenciam ao AG1 e obtiveram pontuação igual ou superior a 8 quando aplicadas as QAQ em seus títulos, resumos e palavras chaves.

Para a avaliação booleana, após a seleção dos artigos finais (AG2), aplicou-se análise qualitativa complementar mediante matriz booleana de avaliação ( $n$  artigos  $\times$  6 QAQ), na qual cada critério QAQ foi avaliado de forma dicotômica (*TRUE/FALSE*) com base na leitura completa dos textos. Esta sistemática booleana, aplicada exclusivamente aos artigos finais, permitiu identificar quais estudos atenderam integralmente cada critério qualitativo específico, gerando score individual de 0 a 6 pontos (formato xx/6) por artigo. A avaliação booleana considerou apenas atendimento completo aos aspectos de cada QAQ (*TRUE* = 1 ponto) ou não atendimento (*FALSE* = 0 pontos), sem graduação intermediária, possibilitando análise objetiva da distribuição de qualidade metodológica na amostra final e identificação de *gaps* críticos nos estudos selecionados.

## 2.7 Estratégia de Síntese e Análise de Dados (síntese analítica multinível)

Esta Revisão Sistemática da Literatura (RSL) adotou uma estratégia de síntese analítica multinível, estruturada em quatro camadas hierárquicas, com o propósito de responder ao objetivo da pesquisa por meio de um processo de agregação progressiva de evidências. A abordagem metodológica teve como referência as diretrizes PRISMA para garantir transparência, reprodutibilidade e rastreabilidade entre dados primários extraídos dos artigos selecionados e as conclusões finais da revisão.

## 2.8 Estrutura Hierárquica das Questões de Pesquisa

A estratégia metodológica foi estruturada em quatro níveis hierárquicos de questões, organizados em forma de cascata, partindo do nível mais abstrato (objetivo geral da RSL) até o nível mais granular (questões de extração de dados). Esta estrutura permitiu o desdobramento sistemático do objetivo em componentes operacionalizáveis, facilitando tanto a coleta quanto a síntese de evidências. A seguir estão apresentados os níveis hierárquicos e estrutura das questões:

Nível 1 - Objetivo da RSL (macro): Define o propósito central da investigação, estabelecendo as dimensões analíticas principais: viabilidade técnica, benefícios e limitações, preservação de privacidade e conformidade regulatória.

Nível 2 - Questões de Pesquisa - QP (intermediário): O objetivo foi desdobrado em três questões de pesquisa (QP1, QP2 e QP3), cada uma abordando dimensões específicas do fenômeno investigado. A QP1 focou na viabilidade do aprendizado federado como alternativa para desenvolvimento de soluções de IA seguras em diagnóstico laboratorial. A QP2 investigou os desafios técnicos, operacionais e regulatórios para adoção da tecnologia. A QP3 explorou comparações entre aprendizado federado e modelos centralizados sob a ótica de privacidade e segurança de dados.

Nível 3 - Questões de Avaliação da Qualidade - QAQ (operacional): As três questões de pesquisa foram subdivididas em seis Questões de Avaliação da Qualidade (QAQ), aplicadas durante a fase de condução para classificar os artigos quanto à sua elegibilidade ao estudo. Cada QP originou duas QAQ complementares, totalizando seis critérios de avaliação: QAQ 1.1 e 1.2 (QP1); QAQ 2.1 e 2.2 (QP2) e QAQ 3.1 e 3.2 (QP3). As QAQ foram aplicadas nos artigos como critério para integrar os Artigos do Grupo 2 (AG2).

Nível 4 - Questões de Extração de Dados - QED (granular): As QAQ foram operacionalizadas por meio de 25 Questões de Extração de Dados (QED), conforme Figura 2, que constituíram o formulário estruturado de extração aplicado aos artigos selecionados. As QED abrangeram aspectos metodológicos (tipo de estudo, população, métricas), técnicos (arquitetura, tecnologias, mecanismos de agregação), regulatórios (conformidade, aspectos éticos) e avaliativos (desafios, benefícios, limitações).

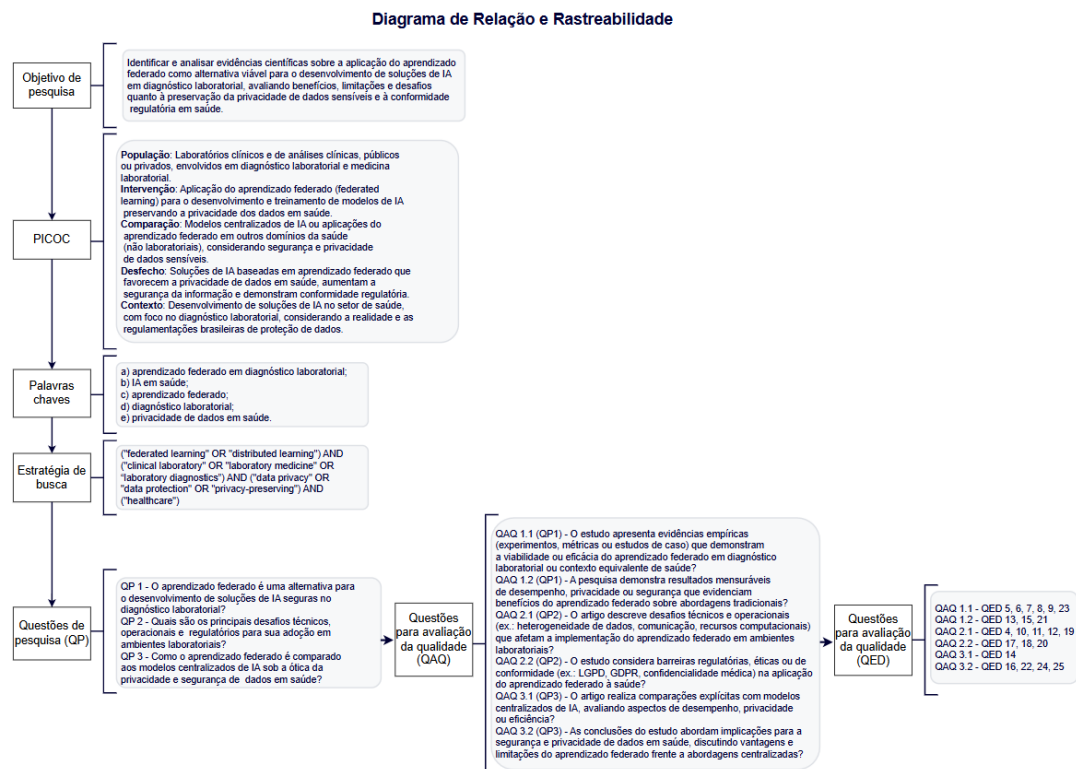
Questões para Extração de Dados (QED)

- 1 - Título do artigo:
- 2 - Ano de publicação:
- 3 - País ou região de estudo:
- 4 - Contexto laboratorial [selecionar...]:
  - a. ☐ ambiente laboratorial hospitalar
  - b. ☐ ambiente laboratorial multicêntrico
  - c. ☐ ambiente laboratorial privado
  - d. ☐ ambiente laboratorial rede pública
- 5 - Tipo de estudo [selecionar...]:
  - a. ☐ Estudo de caso
  - b. ☐ Experimental
  - c. ☐ Prova de conceito
  - d. ☐ Relato técnico
  - e. ☐ Revisão
  - f. ☐ Simulação
- 6 - Objetivo do estudo:
- 7 - Domínio de aplicação laboratorial
- 8 - População/Dados utilizados (PICO-P)
- 9 - Implementação prática de FL [selecionar...]:
  - a. ☐ True
  - b. ☐ False
- 10 - Arquitetura de aprendizado federado:
- 11 - Tecnologias e frameworks:
- 12 - Mecanismo de treinamento e agregação:
- 13 - Métricas de desempenho (O – Outcome):
- 14 - Comparação com modelos centralizados (C – Comparison):
- 15 - Técnicas de preservação da privacidade:
- 16 - Avaliação de segurança e privacidade:
- 17 - Aspectos éticos e conformidade regulatória (LGPD, GDPR):
- 18 - Nível de conformidade regulatória [selecionar...]:
  - a. ☐ Grau de aderência Nenhuma
  - b. ☐ Grau de aderência Parcial
  - c. ☐ Grau de aderência Total
- 19 - Desafios técnicos identificados:
- 20 - Desafios operacionais/regulatórios:
- 21 - Benefícios reportados:
- 22 - Limitações do estudo:
- 23 - Relevância para diagnóstico laboratorial [selecionar...]:
  - a. ☐ Grau de aplicabilidade para laboratórios clínicos Alta
  - b. ☐ Grau de aplicabilidade para laboratórios clínicos Baixa
  - c. ☐ Grau de aplicabilidade para laboratórios clínicos Média
- 24 - Principais conclusões dos autores:
- 25 - Lacunas e recomendações futuras:

**Figura 2. Questões para Extração de Dados (QED)**

## 2.9 Diagrama de Relação e Rastreabilidade

Para assegurar a rastreabilidade entre os níveis hierárquicos, foi construído um diagrama de relação e rastreabilidade que mapeia explicitamente as relações entre QED, QAQ, QP e o objetivo da RSL (Figura 3). Este diagrama estabelece quais QED fundamentam cada QAQ, quais QAQ respondem a cada QP, e como as QP, em conjunto, respondem ao objetivo central. O mapeamento permitiu que cada conclusão da RSL fosse rastreada até os dados primários dos artigos, garantindo transparência e validade interna das inferências.



**Figura 3. Diagrama de relação e rastreabilidade**

## 2.10 - Construção da Matriz Booleana de Avaliação da Qualidade

A aplicação das QAQ aos artigos selecionados resultou na construção de uma matriz booleana de dimensão  $n \times 6$ , onde  $n$  representa o número de artigos selecionados e as seis colunas correspondem às seis QAQ estabelecidas. Para cada artigo, aplicou-se um conjunto de critérios lógicos baseados nos dados extraídos via QED, resultando em valores booleanos (verdadeiro ou falso) para cada QAQ.

Os critérios de avaliação para cada QAQ foram definidos da seguinte forma:

**QAQ 1.1 (Evidências empíricas de viabilidade):** Classificada como verdadeira quando o estudo apresentou delineamento empírico (experimental, estudo de caso ou prova de conceito), reportou implementação prática de aprendizado federado e demonstrou relevância alta ou média para o contexto de diagnóstico laboratorial. A avaliação baseou-se nas QED referentes ao tipo de estudo, implementação prática e relevância clínica.

**QAQ 1.2 (Resultados mensuráveis de desempenho):** Considerada verdadeira para estudos que reportaram métricas quantitativas de desempenho (acurácia, AUC, F1-score, sensibilidade, especificidade, entre outras) e realizaram comparação explícita com abordagens tradicionais ou centralizadas. A presença de métricas e comparações foi verificada por meio das QED específicas para estes aspectos.

**QAQ 2.1 (Desafios técnicos e operacionais):** Avaliada como verdadeira quando o artigo descreveu explicitamente desafios técnicos ou operacionais enfrentados na implementação do aprendizado federado em ambientes laboratoriais, incluindo aspectos como heterogeneidade de dados, comunicação entre nós, recursos computacionais ou

sincronização de modelos. A classificação baseou-se na análise textual das QED relacionadas a desafios técnicos.

QAQ 2.2 (Barreiras regulatórias e conformidade): Classificada como verdadeira para estudos que consideraram aspectos éticos, regulatórios ou de conformidade (LGPD, GDPR, HIPAA ou legislações equivalentes) na aplicação do aprendizado federado à saúde, e que apresentaram nível de conformidade total ou parcial com as regulamentações aplicáveis. A avaliação envolveu as QED de aspectos éticos, conformidade regulatória e desafios operacionais.

QAQ 3.1 (Comparação explícita com modelos centralizados): Considerada verdadeira quando o artigo realizou *benchmarking* direto entre aprendizado federado e modelos centralizados de IA, com apresentação de métricas comparativas de desempenho, privacidade ou eficiência. A verificação foi realizada por meio da QED específica para comparação com modelos centralizados.

QAQ 3.2 (Implicações para segurança e privacidade): Avaliada como verdadeira para estudos cujas conclusões abordaram implicações práticas para segurança e privacidade de dados em saúde, discutindo explicitamente vantagens e limitações do aprendizado federado frente a abordagens centralizadas. A classificação considerou as QED de avaliação de segurança, limitações do estudo e conclusões dos autores.

A matriz booleana resultante serviu como base para as análises subsequentes, permitindo quantificar o número de artigos que atenderam a cada critério de qualidade e identificar padrões de conformidade metodológica no conjunto de estudos selecionados. Adicionalmente, foi calculado um score de qualidade para cada artigo, correspondente ao somatório de QAQ atendidas, variando de 0 a 6, possibilitando análises de distribuição de qualidade metodológica na amostra.

## 2.11 Processo de Síntese Multinível

A síntese dos resultados foi realizada por meio de um processo ascendente, agregando evidências dos níveis mais granulares aos mais abstratos, em quatro etapas sequenciais:

Etapa 1 - Extração granular (QED): Aplicação do formulário estruturado de 25 QED aos artigos selecionados (AG2), com registro padronizado de informações em campos categóricos (múltipla escolha), numéricos (métricas quantitativas) e textuais (descrições qualitativas). Esta etapa resultou em uma base de dados estruturada contendo o total de registros correspondente ao produto entre o número de artigos e o número de QED.

Etapa 2 - Avaliação da qualidade (QED => QAQ): Transformação dos dados extraídos em valores booleanos na matriz de avaliação da qualidade, mediante aplicação dos critérios lógicos estabelecidos para cada QAQ. Esta etapa permitiu identificar, para cada artigo, quais critérios de qualidade foram atendidos, gerando uma matriz que possibilitou análises quantitativas de suficiência metodológica.

Etapa 3 - Síntese das questões de pesquisa (QAQ => QP): Agregação dos resultados das QAQ para responder às questões de pesquisa.

Para a QP1, foram considerados artigos que atenderam simultaneamente às QAQ 1.1 e 1.2, indicando presença tanto de evidências empíricas quanto de resultados mensuráveis.



Para a QP2, foram agregados artigos que atenderam à QAQ 2.1 ou à QAQ 2.2, indicando documentação de desafios técnicos e/ou regulatórios.

Para a QP3, foram considerados estudos que atenderam simultaneamente às QAQ 3.1 e 3.2, demonstrando comparação explícita com análise crítica de implicações.

Esta etapa produziu sínteses narrativas fundamentadas em evidências quantificáveis, identificando padrões, frequências e tendências nos dados agregados.

Etapa 4 - Resposta ao objetivo (QP => Objetivo): Integração das respostas às três questões de pesquisa para construir uma síntese final que abordasse as quatro dimensões do objetivo da RSL: viabilidade técnica, benefícios e limitações, preservação de privacidade e conformidade regulatória.

Esta síntese final foi estruturada de forma a permitir rastreamento completo das conclusões até os dados primários, garantindo que cada afirmação pudesse ser justificada por meio da cadeia QED => QAQ => QP => Objetivo.

## 2.12 Análise de Dados e resultados

As análises tiveram seus resultados representados por estatística descritiva para variáveis categóricas (frequências absolutas e relativas) e síntese narrativa para dados qualitativos. A distribuição de artigos por contexto laboratorial, tipo de estudo, nível de conformidade regulatória e relevância para diagnóstico laboratorial foi caracterizada por meio de tabelas de frequência e representações gráficas. Para as variáveis textuais das QED, aplicou-se análise temática qualitativa com categorização.

A matriz booleana possibilitou análises de cobertura metodológica, permitindo quantificar a proporção de artigos que atenderam a cada critério de qualidade e identificar lacunas na literatura. Adicionalmente, a análise da distribuição de scores de qualidade (0 a 6) permitiu caracterizar a maturidade metodológica do campo investigado.

A análise de coautoria foi conduzida em ambiente RStudio, utilizando a linguagem R versão 4.5.2 (2025-10-31 ucrt) e o pacote bibliometrix [\[Aria and Cuccurullo 2017\]](#), aplicado aos conjuntos de Artigos do Grupo 2 (AG2). A função biblioNetwork() foi utilizada para construir a matriz de colaboração entre autores, a partir da qual se gerou a rede de coautoria. Em seguida, a função networkPlot() foi empregada para visualizar os autores mais relevantes, adotando o *layout* de [Fruchterman and Reingold \(1991\)](#), um algoritmo de desenho de grafos por forças que distribui os nós no plano de forma a equilibrar a repulsão entre vértices e a atração ao longo das arestas, favorecendo *layouts* legíveis e com menor sobreposição de arestas.

## 2.13 Garantia de Validade e Transparência

O processo de síntese multinível implementado nesta RSL assegura três propriedades metodológicas fundamentais: (1) rastreabilidade, por meio do mapeamento explícito entre níveis hierárquicos; (2) transparência, mediante registro estruturado de critérios de avaliação e decisões metodológicas; e (3) reprodutibilidade, possibilitando que pesquisadores independentes possam replicar o processo de síntese a partir dos dados brutos extraídos.

### 3. Resultados

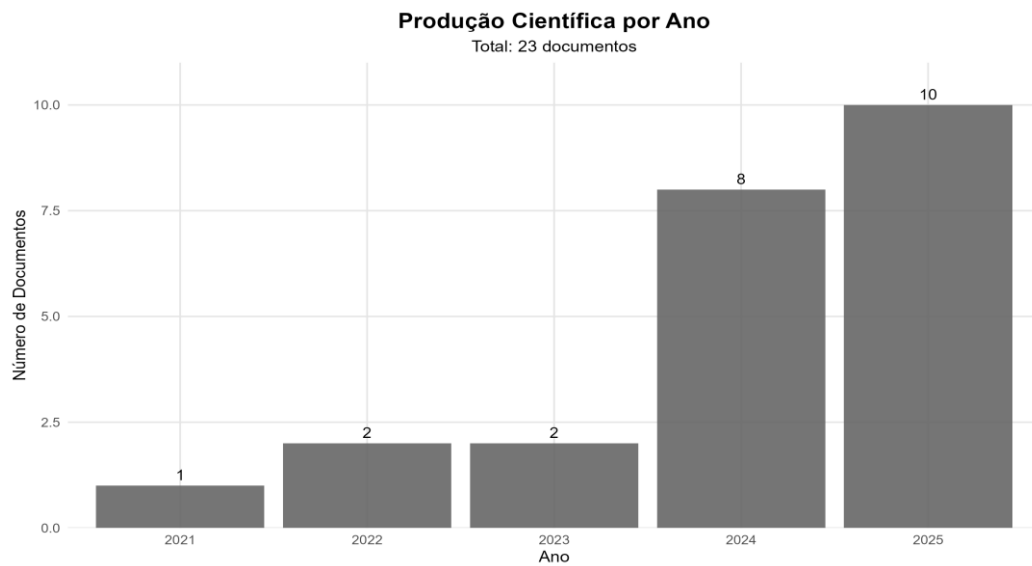
#### 3.1 Caracterização Bibliométrica da Amostra

Após a etapa de seleção dos artigos, restaram os Artigos do Grupo 2 (AG2), constituído por 23 artigos selecionados, sendo os artigos definidos para esta revisão sistemática. Tais artigos possuem seus autores e respectivos anos descritos na Tabela 1.

**Tabela 1. Autores e anos dos artigos selecionados – AG2**

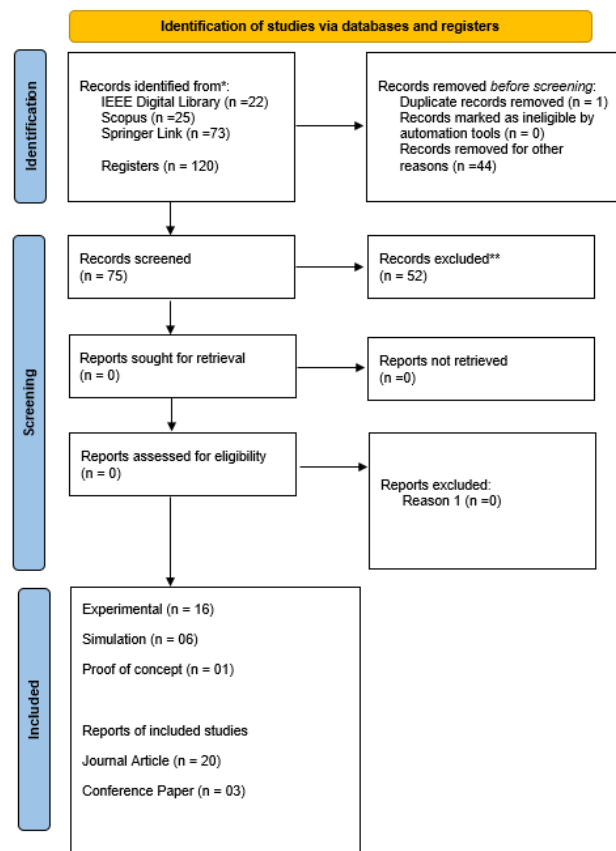
| ID | Autores/ano                                                  |
|----|--------------------------------------------------------------|
| 1  | <a href="#">[Mnasri, S. et al. 2025]</a>                     |
| 2  | <a href="#">[Wei, M. et al. 2025]</a>                        |
| 3  | <a href="#">[Hajder, M. et al. 2021]</a>                     |
| 4  | <a href="#">[Li, Z. et al. 2024]</a>                         |
| 5  | <a href="#">[Gulati, S. et al. 2024]</a>                     |
| 6  | <a href="#">[Fang, C. et al. 2024]</a>                       |
| 7  | <a href="#">[Chen, B. et al. 2024]</a>                       |
| 8  | <a href="#">[Sai, S. et al. 2024]</a>                        |
| 9  | <a href="#">[Lu, M. Y. et al. 2022]</a>                      |
| 10 | <a href="#">[Perales, E. et al. 2025]</a>                    |
| 11 | <a href="#">[Park, E. et al. 2022]</a>                       |
| 12 | <a href="#">[Kumar, D. et al. 2025]</a>                      |
| 13 | <a href="#">[Shahnazeer, C. K. and Sureshkumar, G. 2025]</a> |
| 14 | <a href="#">[Qiu, W. et al. 2024]</a>                        |
| 15 | <a href="#">Zhang 2023</a>                                   |
| 16 | <a href="#">[Düsing, C. et al. 2024]</a>                     |
| 17 | <a href="#">[Liu, X. et al. 2025]</a>                        |
| 18 | <a href="#">[He, P. et al. 2023]</a>                         |
| 19 | <a href="#">[Zhang, H. et al. 2025]</a>                      |
| 20 | <a href="#">[Zou, W. et al. 2025]</a>                        |
| 21 | <a href="#">[Baid, U. et al. 2024]</a>                       |
| 22 | <a href="#">[Ma, T. et al. 2025]</a>                         |
| 23 | <a href="#">[Jamshidi, M. et al. 2025]</a>                   |

Todos foram publicados entre 2021 e 2025, evidenciando a contemporaneidade da produção científica sobre aprendizado federado em diagnóstico em saúde e laboratorial. Conforme demonstrado na Figura 4, a análise temporal revela crescimento expressivo no período, com pico de produção em 2025, quando foram publicados 10 documentos (43,5% do total), seguido por 8 artigos em 2024 (34,8%). Os anos de 2022 e 2023 apresentaram produção incipiente, com 2 artigos em cada ano (8,7% cada), enquanto 2021 registrou apenas 1 artigo (4,3%), indicando que o campo começou a consolidar-se a partir de 2024.



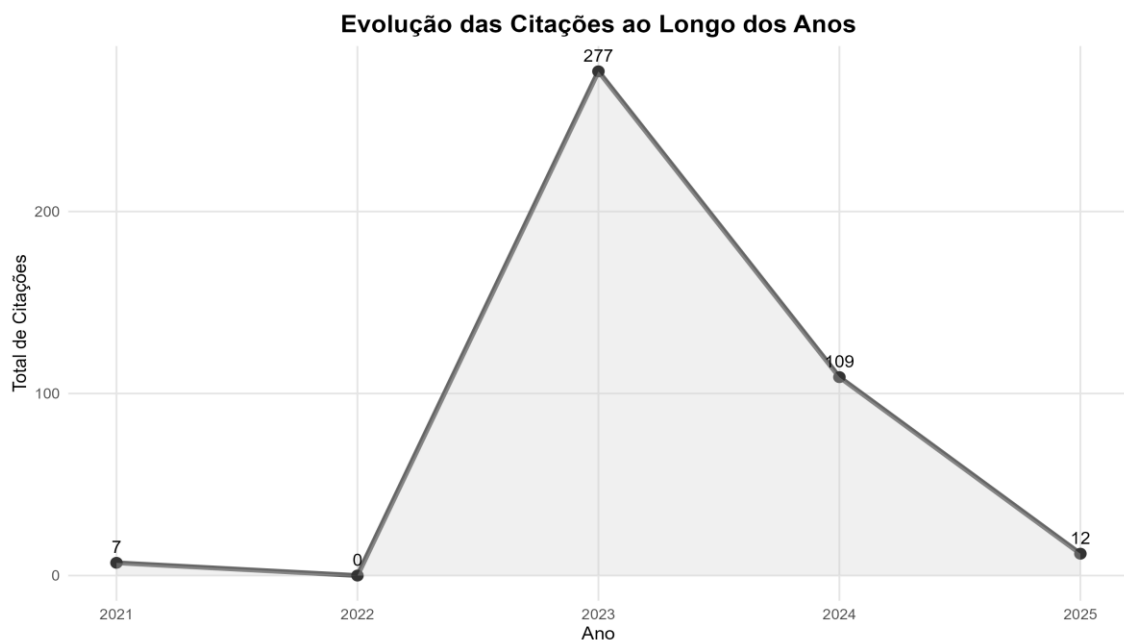
**Figura 4. Produção científica por ano (AG2, n = 23)**

Como referência, foi seguida a abordagem Prisma [Page et al. 2021], dessa forma os resultados podem ser observados no diagrama da Figura 5.



**Figura 2. Diagrama de fluxo Prisma**

A distribuição de citações revela padrão temporal concentrado, com pico acentuado em 2023, quando os artigos acumularam 277 citações, representando 68,4% do total de citações da amostra. Este resultado evidencia alto impacto inicial de trabalhos publicados nesse ano, que apresentaram média de 138,5 citações por artigo (2 artigos publicados). Em 2024, os 8 artigos publicados acumularam 109 citações, com média de 13,6 citações por documento. O total de 405 citações dos 23 artigos resulta em média geral de 17,6 citações por documento, demonstrando relevância científica da amostra selecionada. Os artigos de 2021 registraram 7 citações, enquanto os de 2025 acumularam 12 citações até o momento da análise. Os trabalhos de 2022 não receberam citações, conforme pode ser constatado na Figura 6. A tendência decrescente observada após 2023 é esperada, considerando o tempo decorrido desde a publicação, sendo que artigos mais recentes (2024 e 2025) ainda estão em fase de consolidação de seu impacto científico.



**Figura 6 Evolução das citações ao longo dos anos**

A análise de autoria revela participação de 124 autores únicos distribuídos pelos 23 artigos, indicando média de aproximadamente 5,4 autores por publicação. Não houve concentração expressiva de produção em autores específicos, uma vez que a listagem dos 15 autores mais frequentes demonstra que cada um contribuiu com apenas 1 artigo na amostra final. Essa dispersão autoral sugere campo de pesquisa ainda em formação, sem grupos consolidados dominantes, característica típica de áreas emergentes.

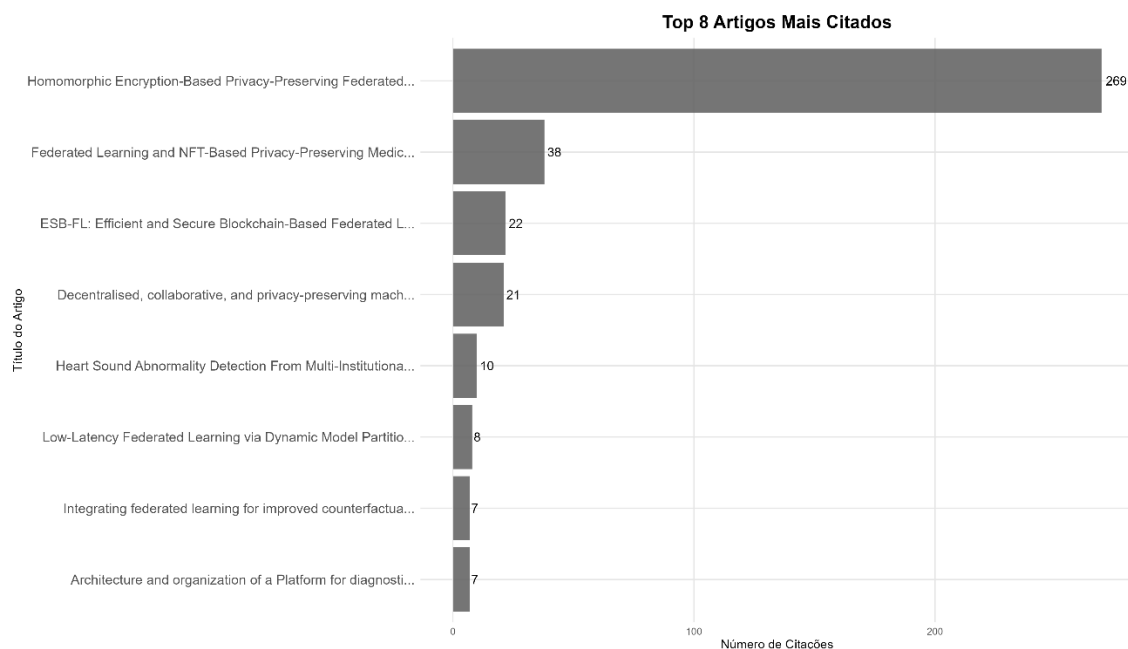
O conjunto de 109 palavras-chave únicas extraídas dos artigos evidencia diversidade temática dentro do escopo da revisão. A palavra-chave "*federated learning*" aparece em 15 dos 23 artigos (65,2%) confirmando centralidade do tema. As palavras "*privacy protection*" (3 ocorrências), "*artificial intelligence*" (2 ocorrências), "*blockchain*" (2 ocorrências), "*convolutional neural networks*" (2 ocorrências), "*smart contract*" (2 ocorrências) e "*split learning*" (2 ocorrências) destacam-se como termos secundários recorrentes, revelando associação frequente do aprendizado federado com

tecnologias de proteção de privacidade (*blockchain, smart contracts*) e arquiteturas de *deep learning*.

A distribuição por periódicos demonstra pulverização da produção científica em múltiplos veículos de publicação. O IEEE Internet of Things Journal lidera com 3 artigos, seguido por IEEE Journal of Biomedical and Health Informatics e Procedia Computer Science, ambos com 2 artigos cada. Os demais 16 artigos distribuem-se por diferentes periódicos e conferências, cada um publicando 1 trabalho da amostra. As publicações do IEEE concentram aproximadamente 52% da amostra (12 de 23 artigos), evidenciando predominância de veículos voltados a engenharia, computação e sistemas de saúde.

A análise de impacto por periódico revela que o IEEE Internet of Things Journal, com 3 artigos, acumula 42 citações (média de 14,0 citações por artigo). O IEEE Transactions on Big Data, com 1 artigo, apresenta 22 citações. Esses dados indicam que trabalhos publicados em periódicos de alto impacto, mesmo em menor quantidade, exercem influência significativa sobre o campo.

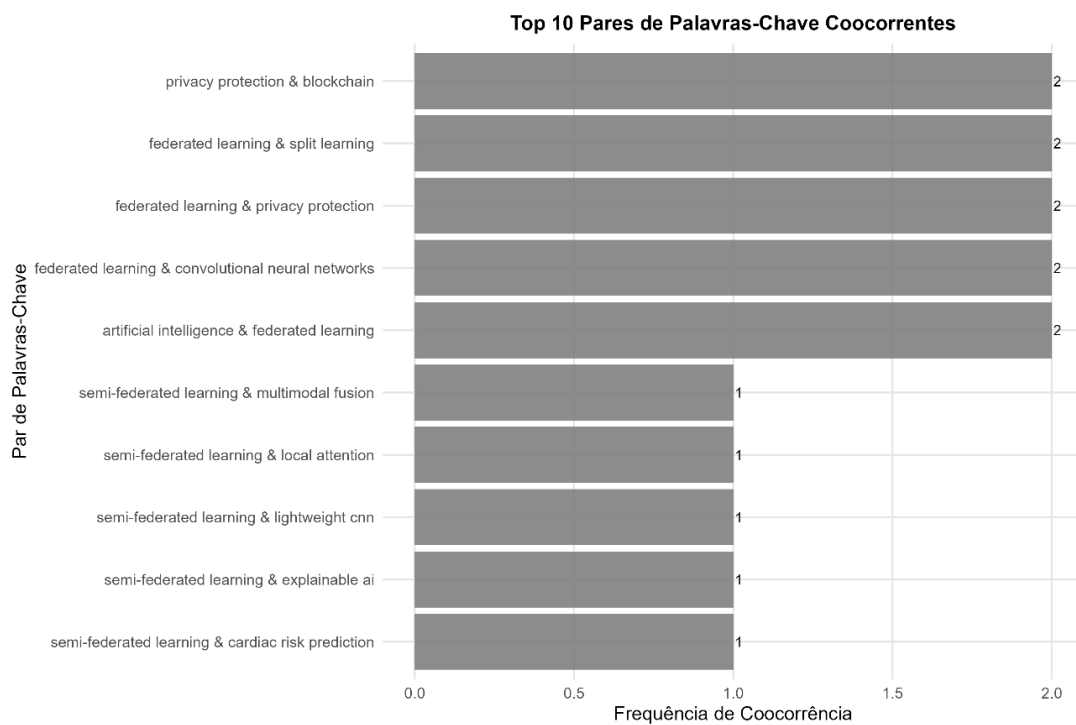
O artigo mais citado da amostra, "*Homomorphic Encryption-Based Privacy-Preserving Federated Learning in IoT-Enabled Healthcare System*", publicado no *IEEE Transactions on Network Science and Engineering* em 2023, acumula 269 citações, representando 66,4% do total de citações da amostra (Figura 7). Esse trabalho, seguido por outros sobre blockchain, NFT e arquiteturas descentralizadas, evidencia forte interesse da comunidade científica em soluções que combinam aprendizado federado com tecnologias de criptografia e arquiteturas distribuídas para preservação de privacidade.



**Figura 3. Artigos mais citados (top 8 artigos)**

A análise de pares de palavras-chave coocorrentes identifica cinco agrupamentos principais com frequência igual ou superior a 2 ocorrências: "*privacy protection & blockchain*" (2 ocorrências), "*federated learning & split learning*" (2 ocorrências),

"*federated learning & privacy protection*" (2 ocorrências), "*federated learning & convolutional neural networks*" (2 ocorrências) e "*artificial intelligence & federated learning*" (2 ocorrências). Esses pares refletem tendências de pesquisa que associam aprendizado federado a tecnologias específicas de proteção de dados e arquiteturas de redes neurais profundas como demonstrado na Figura 8.



**Figura 4. Pares de palavras chaves (top 10 pares)**

### 3.2 Avaliação da Qualidade Metodológica

A aplicação das seis Questões de Avaliação da Qualidade (QAQ) aos 23 artigos selecionados resultou em matriz booleana que permitiu caracterizar a suficiência metodológica da amostra em relação aos critérios estabelecidos. A análise revela padrões de cobertura heterogêneos entre as diferentes dimensões de qualidade avaliadas.

A QAQ 1.1, que avalia presença de evidências empíricas de viabilidade do aprendizado federado em contextos de diagnóstico laboratorial ou saúde, foi atendida por 13 artigos (56,5% da amostra). Esse resultado indica que mais da metade dos estudos selecionados apresentou delineamento experimental, estudo de caso ou prova de conceito com implementação prática de aprendizado federado aplicada a contextos com relevância alta ou média para diagnóstico laboratorial. Os 10 artigos que não atenderam esse critério (43,5%) correspondem predominantemente a simulações realizadas em ambientes virtuais, propostas arquiteturais sem validação empírica ou estudos cujo foco primário não foi aplicação laboratorial direta.

A QAQ 1.2, que examina presença de resultados mensuráveis de desempenho demonstrando benefícios do aprendizado federado, também foi atendida por 13 artigos (56,5%). Observou-se correlação perfeita entre QAQ 1.1 e QAQ 1.2: todos os artigos que

apresentaram evidências empíricas de viabilidade também reportaram métricas quantitativas de desempenho. Essa correlação sugere que estudos com implementação prática tendem naturalmente a produzir avaliações quantitativas de eficácia, enquanto trabalhos puramente teóricos ou arquiteturais não alcançam essa etapa de validação.

A QAQ 2.1, referente à descrição de desafios técnicos e operacionais na implementação do aprendizado federado, foi atendida por todos os 23 artigos (100%). Esse resultado indica que a documentação de barreiras técnicas (heterogeneidade de dados, custos de comunicação, sincronização de modelos, recursos computacionais) constitui elemento universal na literatura analisada, independentemente do tipo de estudo ou nível de maturidade da implementação. Mesmo trabalhos teóricos ou de simulação dedicaram atenção explícita aos desafios técnicos inerentes ao paradigma federado.

A QAQ 2.2, que avalia consideração de barreiras regulatórias, éticas ou de conformidade, foi atendida por apenas 6 artigos (26,1%), representando o critério de qualidade com menor cobertura na amostra. Esse achado revela lacuna crítica: enquanto a totalidade dos estudos aborda desafios técnicos, apenas um quarto menciona explicitamente aspectos de conformidade com legislações de proteção de dados (LGPD, GDPR, HIPAA) ou discute implicações éticas da aplicação do aprendizado federado em saúde. Os 17 artigos restantes (73,9%) não fizeram menção direta a frameworks regulatórios ou limitaram-se a citações genéricas sobre privacidade sem conexão explícita com arcabouços legais aplicáveis.

A QAQ 3.1, referente à comparação explícita com modelos centralizados de IA, foi atendida por 15 artigos (65,2%). Esse resultado demonstra que aproximadamente dois terços da amostra realizaram *benchmarking* direto entre aprendizado federado e abordagens tradicionais centralizadas, apresentando métricas comparativas de desempenho, privacidade ou eficiência. Os 8 artigos que não atenderam esse critério (34,8%) focaram exclusivamente na avaliação do aprendizado federado sem estabelecer linha de base comparativa com modelos centralizados.

A QAQ 3.2, que examina se as conclusões dos autores abordam implicações para segurança e privacidade de dados em saúde com discussão de vantagens e limitações, foi atendida por 9 artigos (39,1%). Esse resultado indica que menos da metade dos estudos realizou avaliação formal de segurança, explicitou limitações metodológicas e discutiu *trade-offs* entre privacidade, desempenho e viabilidade prática. Os 14 artigos que não atenderam esse critério (60,9%) apresentaram conclusões focadas em aspectos técnicos de desempenho sem aprofundamento crítico sobre implicações de segurança ou sem reconhecimento explícito de limitações da abordagem proposta.

A matriz booleana está representada na Tabela 2, onde estão os resultados das análises referentes às QAQ para cada artigo do AG2, incluindo coluna com os valores do total de atendimento às QAQ alcançado por artigo.

Tabela 2. Tabela booleana QAQ

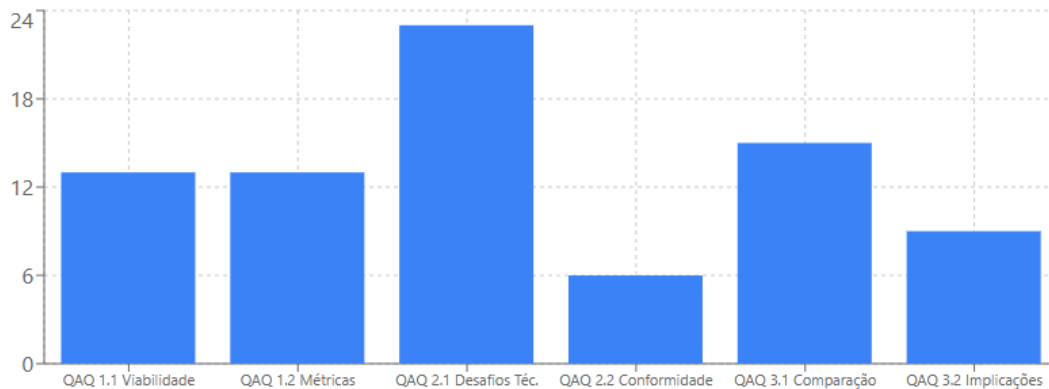
| ID | Título                                                    | Autores/ano                | QAQ 1.1 | QAQ 1.2 | QAQ 2.1 | QAQ 2.2 | QAQ 3.1 | QAQ 3.2 | Total | TRUE | Percentual | TRUE   |
|----|-----------------------------------------------------------|----------------------------|---------|---------|---------|---------|---------|---------|-------|------|------------|--------|
| 1  | A Hybrid Blockchain and Federated Learning ...            | [Mnasri, S. et al. 2025]   | TRUE    | TRUE    | TRUE    | FALSE   | TRUE    | TRUE    | 5     |      |            | 83,33  |
| 2  | AccDFL: Accelerated Decentralized Federated ...           | [Wei, M. et al. 2025]      | TRUE    | TRUE    | TRUE    | FALSE   | TRUE    | TRUE    | 5     |      |            | 83,33  |
| 3  | Architecture and organization of a Platform for ...       | [Hajder, M. et al. 2021]   | FALSE   | FALSE   | TRUE    | TRUE    | FALSE   | FALSE   | 2     |      |            | 33,33  |
| 4  | Blockchain-based collaborative data analysis ...          | [Li, Z. et al. 2024]       | TRUE    | TRUE    | TRUE    | FALSE   | TRUE    | FALSE   | 4     |      |            | 66,67  |
| 5  | Collaborative Privacy-Preserving Federated Learning ...   | [Gulati, S. et al. 2024]   | TRUE    | TRUE    | TRUE    | FALSE   | TRUE    | FALSE   | 4     |      |            | 66,67  |
| 6  | Decentralised collaborative and privacy-preserving ...    | [Fang, C. et al. 2024]     | TRUE    | TRUE    | TRUE    | FALSE   | TRUE    | TRUE    | 5     |      |            | 83,33  |
| 7  | ESB-FL: Efficient and Secure Blockchain-Based ...         | [Chen, B. et al. 2024]     | TRUE    | TRUE    | TRUE    | FALSE   | FALSE   | TRUE    | 4     |      |            | 66,67  |
| 8  | Federated Learning and NFT-Based Privacy-Preserving ...   | [Sai, S. et al. 2024]      | FALSE   | FALSE   | TRUE    | FALSE   | FALSE   | FALSE   | 1     |      |            | 16,67  |
| 9  | Federated learning for computational pathology ...        | [Lu, M. Y. et al. 2022]    | TRUE    | TRUE    | TRUE    | FALSE   | TRUE    | TRUE    | 5     |      |            | 83,33  |
| 10 | Federated Learning for Securing Medical Imaging ...       | [Perales, E. et al. 2025]  | TRUE    | TRUE    | TRUE    | TRUE    | TRUE    | TRUE    | 6     |      |            | 100,00 |
| 11 | Federated Learning Models using Flow Cytometry ...        | [Park, E. et al. 2022]     | FALSE   | FALSE   | TRUE    | FALSE   | FALSE   | FALSE   | 1     |      |            | 16,67  |
| 12 | Federated learning with explainable AI for liver ...      | [Kumar, D. et al. 2025]    | TRUE    | TRUE    | TRUE    | TRUE    | TRUE    | TRUE    | 6     |      |            | 100,00 |
| 13 | Federated Transfer Learning Framework for ...             | [Shahnazeer, C. K. 2025]   | FALSE   | FALSE   | TRUE    | FALSE   | FALSE   | FALSE   | 1     |      |            | 16,67  |
| 14 | Heart Sound Abnormality Detection From ...                | [Qiu, W. et al. 2024]      | TRUE    | TRUE    | TRUE    | FALSE   | TRUE    | TRUE    | 5     |      |            | 83,33  |
| 15 | Homomorphic Encryption-Based Privacy-Preserving ...       | [Zhang, L. et al. 2023]    | TRUE    | TRUE    | TRUE    | FALSE   | FALSE   | TRUE    | 4     |      |            | 66,67  |
| 16 | Integrating federated learning for improved ...           | [Düsing, C. et al. 2024]   | FALSE   | FALSE   | TRUE    | FALSE   | TRUE    | FALSE   | 2     |      |            | 33,33  |
| 17 | Interpretable Semi-federated Learning for Multimodal ...  | [Liu, X. et al. 2025]      | FALSE   | FALSE   | TRUE    | TRUE    | TRUE    | FALSE   | 3     |      |            | 50,00  |
| 18 | Low-Latency Federated Learning via Dynamic Model ...      | [He, P. et al. 2023]       | FALSE   | FALSE   | TRUE    | FALSE   | FALSE   | FALSE   | 1     |      |            | 16,67  |
| 19 | Non-JID Medical Image Segmentation Based ...              | [Zhang, H. et al. 2025]    | FALSE   | FALSE   | TRUE    | FALSE   | TRUE    | FALSE   | 2     |      |            | 33,33  |
| 20 | On a Federated-Learning-Based Computation ...             | [Zou, W. et al. 2025]      | FALSE   | FALSE   | TRUE    | FALSE   | FALSE   | FALSE   | 1     |      |            | 16,67  |
| 21 | Pan-Cancer Tumor Infiltrating Lymphocyte ...              | [Baid, U. et al. 2024]     | FALSE   | FALSE   | TRUE    | TRUE    | TRUE    | FALSE   | 3     |      |            | 50,00  |
| 22 | Privacy and Fairness-Guaranteed Federated ...             | [Ma, T. et al. 2025]       | TRUE    | TRUE    | TRUE    | FALSE   | TRUE    | FALSE   | 4     |      |            | 66,67  |
| 23 | Revolutionizing biological digital twins: Integrating ... | [Jamshidi, M. et al. 2025] | TRUE    | TRUE    | TRUE    | TRUE    | TRUE    | FALSE   | 5     |      |            | 83,33  |

A distribuição dos artigos segundo número total de QAQ atendidas revela estratificação da qualidade metodológica. Dois artigos (8,7%) atenderam aos seis critérios de qualidade estabelecidos, demonstrando excelência metodológica em todas as dimensões avaliadas. Seis artigos (26,1%) atenderam cinco critérios, cinco artigos (21,7%) atenderam quatro critérios, quatro artigos (17,4%) atenderam três critérios, um artigo (4,3%) atendeu dois critérios e cinco artigos (21,7%) atenderam apenas um critério. Essa distribuição evidencia heterogeneidade metodológica substancial na literatura, com concentração de aproximadamente 35% da amostra (8 artigos) apresentando alta qualidade metodológica (5 ou 6 critérios atendidos).

Os dois artigos com pontuação máxima (6 critérios atendidos) foram "*Federated Learning for Securing Medical Imaging Against Deepfakes and Adversarial Attacks*" [Perales, E. et al. 2025] e "*Federated learning with explainable AI for liver disease prediction using clinical data*" [Kumar, D. et al. 2025], ambos publicados em 2025. Esses trabalhos destacam-se por combinar evidências empíricas robustas, métricas quantitativas de desempenho, comparação explícita com modelos centralizados, discussão de desafios técnicos, consideração de aspectos regulatórios e análise crítica de implicações para segurança e privacidade. Representam, portanto, referências metodológicas para o campo.

Como demonstrado na Figura 9, a análise entre as QAQ revela padrão de alta coerência entre QAQ 1.1 e QAQ 1.2, como anteriormente mencionado, indicando que a presença de implementação prática prediz fortemente a presença de métricas quantitativas. Por outro lado, a baixa cobertura da QAQ 2.2 (26,1%) combinada com cobertura universal da QAQ 2.1 (100%) sugere dissociação entre preocupações técnicas e preocupações regulatórias na literatura: enquanto desafios técnicos são universalmente documentados, aspectos de conformidade legal permanecem sub representados.



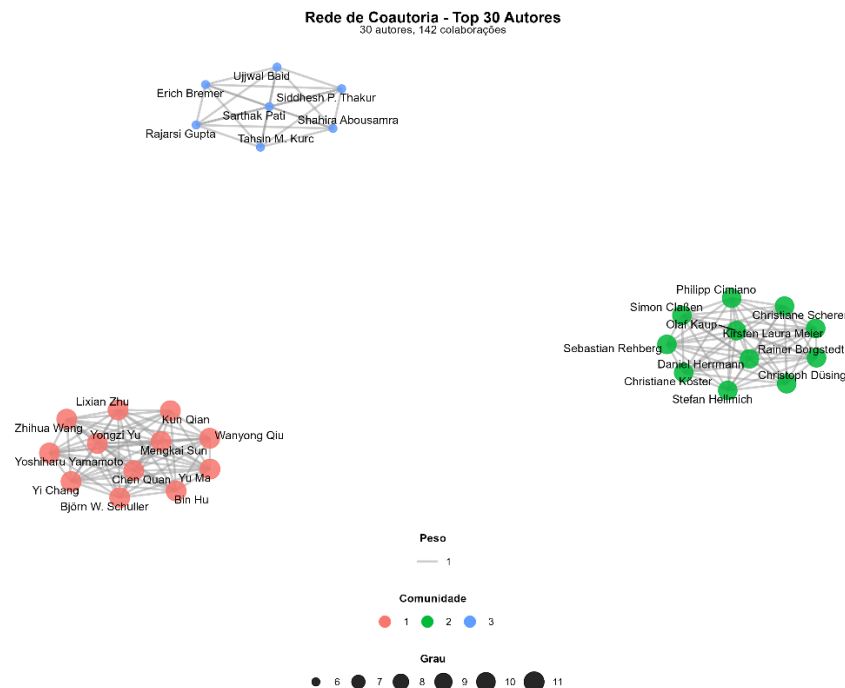


**Figura 5. Atendimento aos critérios QAQ**

### 3.3 Análise de Coautoria e Colaboração Científica

A análise de rede de coautoria foi conduzida utilizando a linguagem R versão 4.5.2 e o pacote bibliometrix [Aria and Cuccurullo 2017], aplicado ao conjunto de 23 artigos que compuseram a amostra final desta revisão sistemática. A função biblioNetwork() foi utilizada para construir a matriz de colaboração entre autores, a partir da qual se gerou a rede de coautoria. Em seguida, a função networkPlot() foi empregada para visualizar os 30 autores mais relevantes, adotando o *layout* de Fruchterman and Reingold (1991), um algoritmo de desenho de grafos por forças que distribui os nós no plano de forma a equilibrar a repulsão entre vértices e a atração ao longo das arestas, favorecendo *layouts* legíveis e com menor sobreposição de arestas.

A rede resultante apresentou 30 autores com 142 colaborações, densidade de 0,3264 e grau médio de 9,47. O algoritmo walktrap identificou três componentes conectados com modularidade de 0,32, valor moderado que indica estrutura de rede parcialmente fragmentada. Essa modularidade intermediária sugere que a rede apresenta alguma segmentação em clusters, mas ainda mantém conexões inter-grupos, caracterizando campo de pesquisa em consolidação com núcleos especializados emergentes, mas não completamente isolados. A análise visual da rede revela três clusters principais com características distintas. Conforme demonstrado na Figura 10, o primeiro cluster, representado em vermelho, organizou-se em torno de Lixian Zhu, que apresenta o maior grau de conectividade, juntamente com colaboradores como Zhihua Wang, Yu Ma, Chen Quan, Bin Hu e Yongzi Yu, apresentando alta densidade de conexões internas e posicionamento central na rede. O segundo cluster, em verde, estruturou-se em torno de Rainer Borgstedt e Daniel Herrmann, com colaboradores europeus incluindo Christiane Scherer, Christiane Köster, Christoph Düsing e Olaf Kaup, caracterizando-se por colaboração intensa e especialização temática. O terceiro cluster, em azul, agregou Shahira Abousamra, Ujjwal Baid, Sarthak Pati, Siddhesh P. Thakur e colaboradores, configurando grupo com padrão de coautoria recorrente associado a projeto de pesquisa multicêntrico específico.



**Figura 10. Rede de coautoria – top 30 autores**

A análise de coautoria revela que o campo do aprendizado federado em diagnóstico laboratorial caracteriza-se por estrutura de colaboração em desenvolvimento, com três núcleos especializados distintos identificados. A modularidade moderada (0,32) e a densidade relativamente baixa (0,3264) indicam que, embora existam grupos colaborativos definidos, o campo ainda não apresenta segmentação rígida, sugerindo potencial para expansão de colaborações intergrupos. O grau médio de 9,47 colaborações por autor evidencia ambiente colaborativo ativo, porém ainda em fase de consolidação de redes estáveis de pesquisa. A presença de três clusters bem delimitados, porém com alguma conectividade entre si, sugere que o campo se encontra em estágio intermediário de maturação científica, com grupos temáticos emergentes, mas sem isolamento completo, característica típica de áreas de pesquisa em expansão.

#### 4. Discussão

A análise sistemática de 23 artigos selecionados revelou um panorama heterogêneo e em evolução quanto à maturidade metodológica das aplicações de aprendizado federado em diagnóstico laboratorial. O score médio de qualidade de 3,43/6, obtido pela matriz booleana de avaliação (23 artigos  $\times$  6 critérios QAQ), indica que o campo apresenta desenvolvimento metodológico intermediário, com substancial variabilidade entre os estudos analisados.

##### 4.1 - Viabilidade Técnica do Aprendizado Federado em Diagnóstico Laboratorial

Os resultados evidenciaram que 13 artigos (56,5%) demonstraram viabilidade técnica do aprendizado federado mediante apresentação de evidências empíricas concretas e

métricas quantitativas de desempenho (QP1). Este achado corrobora a hipótese de que o AF constitui alternativa tecnicamente viável para desenvolvimento de soluções de IA em diagnóstico laboratorial. A correlação perfeita observada entre QAQ 1.1 e QAQ 1.2, sugere maturidade metodológica neste subconjunto de pesquisas experimentais, já que todos os estudos que reportaram evidências empíricas também apresentaram métricas quantitativas. Entretanto, 43,5% dos estudos não atenderam aos critérios de viabilidade técnica, frequentemente por se basearem em simulações computacionais ou propostas conceituais sem validação empírica em ambientes laboratoriais reais. Esta limitação metodológica reduz a generalidade dos achados e aponta para necessidade de estudos experimentais em cenários de saúde reais.

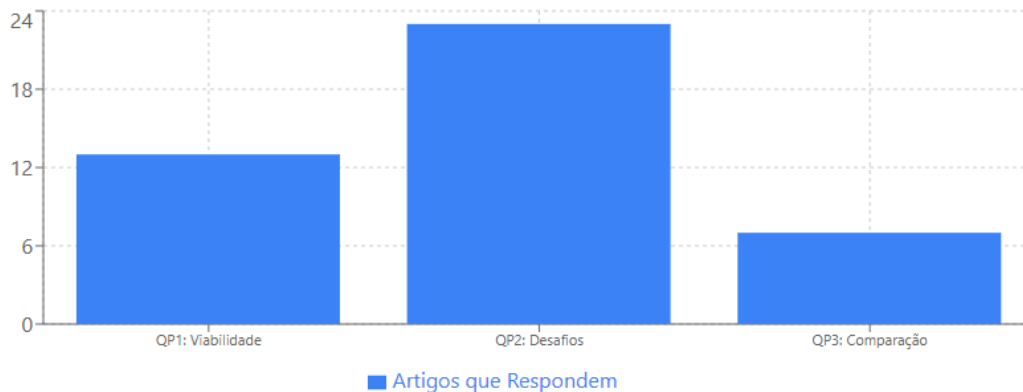
#### 4.2 - Desafios Técnicos e Operacionais

A totalidade dos 23 artigos (100%) documentou desafios técnicos, operacionais ou de implementação relacionados ao aprendizado federado (QP2, via QAQ 2.1). Este consenso absoluto demonstra que a comunidade científica reconhece as barreiras práticas para adoção do AF, incluindo heterogeneidade de dados entre instituições, custos computacionais elevados, complexidade de coordenação entre múltiplos sites e dificuldades na sincronização de modelos distribuídos. O reconhecimento universal de desafios técnicos contrasta fortemente com a baixa frequência de estudos que abordam conformidade regulatória explicitamente. Enquanto 100% dos artigos mencionam barreiras técnicas, apenas 6 estudos (26,1%) fazem referência explícita a *frameworks* regulatórios como GDPR, LGPD ou HIPAA (QAQ 2.2). Este desequilíbrio configura *gap* crítico na literatura, uma vez que conformidade regulatória constitui requisito obrigatório para adoção de tecnologias de IA em ambientes de saúde.

#### 4.3 - Privacidade e Comparação com Modelos Centralizados

A análise revelou que 15 artigos (65,2%) realizaram comparação entre aprendizado federado e modelos centralizados (QAQ 3.1), mas apenas 9 estudos (39,1%) conduziram avaliação formal de segurança com análise crítica de *trade-offs*, limitações e implicações práticas (QP3). Esta diferença indica que, embora comparações técnicas sejam relativamente frequentes, avaliações aprofundadas sobre garantias de privacidade e segurança permanecem insuficientes. A proporção de 39,1% de estudos com avaliação formal completa sugere que o segmento ainda carece de padrões metodológicos robustos para mensuração de efetividade do AF em preservação de privacidade. Muitos estudos apresentam vantagens teóricas do AF sem validação empírica rigorosa ou análise crítica de vulnerabilidades potenciais, como ataques de inferência de participação ou reconstrução de dados sensíveis a partir de gradientes compartilhados.

Na Figura 11 são apresentadas as QP1, QP2 e QP3 e o quantitativo de artigos que responderam cada uma das questões de pesquisa.



**Figura 6. Artigos que responderam as QP**

#### 4.4 - Gap Crítico em Conformidade Regulatória

O achado mais evidente desta revisão é a baixa cobertura de aspectos regulatórios: apenas 26,1% dos estudos (6/23) abordam explicitamente conformidade com legislações de proteção de dados. Considerando que GDPR (União Europeia), LGPD (Brasil) e HIPAA (Estados Unidos) impõem requisitos estritos para processamento de dados de saúde, a negligência deste aspecto em 73,9% dos artigos representa barreira significativa para tradução de pesquisas em implementações clínicas reais.

Este *gap* regulatório pode ser explicado por três fatores: predominância de estudos conduzidos em ambientes acadêmicos ou de pesquisa, sem pressão regulatória imediata; foco técnico dos pesquisadores em otimização de algoritmos, posicionando aspectos legais a segundo plano; ausência de *frameworks* regulatórios específicos para aprendizado federado, gerando incerteza sobre requisitos aplicáveis. A evolução do campo exigirá abordagem interdisciplinar que integre expertise técnica, jurídica e regulatória.

#### 4.5 - Heterogeneidade Metodológica e Qualidade dos Estudos

A distribuição de scores de qualidade evidenciou heterogeneidade metodológica importante. Enquanto 2 artigos (8,69%) atingiram pontuação máxima de 6/6, demonstrando excelência metodológica abrangente, outros estudos apresentaram scores significativamente inferiores, com 8 artigos (34,78%) obtendo apenas 2 pontos ou menos. Esta variabilidade reflete diferentes estágios de maturidade das pesquisas, desde propostas conceituais preliminares até estudos experimentais robustos com validação clínica.

O *score* médio de 3,43/6 indica que o campo se encontra em fase intermediária de desenvolvimento científico. A ausência de padrões metodológicos consolidados e de protocolos de avaliação uniformes contribui para esta heterogeneidade, dificultando comparações diretas entre estudos e síntese de evidências.

#### 4.6 - Limitações da Revisão Sistemática

Esta RSL apresenta limitações que devem ser consideradas na interpretação dos achados. Primeiro aspecto, a amostra de 23 artigos, embora selecionada rigorosamente mediante critérios baseados nos princípios do PRISMA, representa recorte temporal específico

(2021-2025) e pode não capturar toda diversidade de aplicações de AF em diagnóstico laboratorial. Segundo ponto, a avaliação qualitativa mediante matriz booleana, apesar de objetiva e replicável, pode não capturar nuances metodológicas que escalas graduadas identificariam. Terceiro aspecto, a heterogeneidade conceitual do termo "diagnóstico laboratorial" nos estudos analisados, abrangendo desde análises clínicas tradicionais até patologia computacional e diagnóstico por imagem, pode ter introduzido variabilidade não controlada. Quarto aspecto, publicações em idiomas não-inglês ou estudos técnicos não indexados nas bases consultadas podem ter sido inadvertidamente excluídos e gerado viés (*language bias*). Tal limitação deve ser considerada na interpretação dos resultados e aponta para a necessidade de revisões futuras com critérios idiomáticos mais abrangentes. Por fim, o quinto aspecto aponta para o fato de que a avaliação de conformidade regulatória baseou-se exclusivamente em menções explícitas no texto, podendo subestimar estudos que implicitamente seguiram boas práticas regulatórias sem documentação formal.

## 5. Conclusão

Em resposta às questões de pesquisa e ao objetivo proposto, seguem as considerações e conclusões.

Para a QP1 o aprendizado federado constitui alternativa viável para desenvolvimento de soluções de IA seguras em diagnóstico laboratorial, com 56,5% dos estudos (13/23) apresentando evidências empíricas concretas e métricas quantitativas de desempenho.

Para a QP2 os principais desafios técnicos e operacionais foram reconhecidos universalmente (100% dos estudos), mas a conformidade regulatória explícita com GDPR/LGPD/HIPAA foi identificada em apenas 26,1% dos artigos (6/23), configurando gap crítico;

Para a QP3 a comparação com modelos centralizados foi realizada em 65,2% dos estudos (15/23), porém apenas 39,1% (9/23) conduziram avaliação formal de segurança com análise crítica de *trade-offs* entre privacidade, desempenho e viabilidade operacional.

Dessa forma conclui-se que o objetivo da pesquisa foi parcialmente alcançado. A análise sistemática de 23 artigos possibilitou identificar evidências científicas sobre viabilidade técnica do aprendizado federado em diagnóstico laboratorial e mapear desafios técnicos de forma abrangente. Entretanto, limitações metodológicas substanciais foram identificadas: heterogeneidade qualitativa significativa (score médio 3,43/6), insuficiência de avaliações formais de privacidade e segurança, e, especialmente, a não consideração de aspectos regulatórios em 73,9% dos estudos analisados.

Conclui-se que, embora o aprendizado federado demonstre potencial técnico comprovado para diagnóstico laboratorial, sua adoção em larga escala permanece limitada pela ausência de *frameworks* regulatórios específicos, padrões metodológicos consolidados e evidências robustas sobre garantias de privacidade em ambientes clínicos reais. A disparidade entre maturidade técnica (100% reconhecem desafios) e maturidade regulatória (26,1% abordam conformidade) representa a principal barreira para tradução de pesquisas em implementações práticas em ambientes de saúde regulados.

## 6. Recomendações para pesquisas futuras

Com base nos *gaps* identificados, recomendam-se três direções prioritárias para pesquisas futuras:

Desenvolvimento de *frameworks* regulatórios específicos para AF em saúde que leve a pesquisas interdisciplinares envolvendo especialistas em direito digital, proteção de dados e saúde pública, essenciais para traduzir princípios gerais de GDPR/LGPD/HIPAA em diretrizes operacionais aplicáveis ao aprendizado federado em diagnóstico laboratorial.

Padronização metodológica e protocolos de avaliação de qualidade que leve ao estabelecimento de *guidelines* metodológicos específicos para estudos de AF em contextos laboratoriais, incluindo protocolos de avaliação formal de segurança, métricas padronizadas de privacidade e requisitos mínimos para validação empírica em ambientes clínicos reais.

Estudos experimentais multicêntricos em ambientes regulados para condução de ensaios clínicos prospectivos envolvendo múltiplas instituições em contextos regulamentados, com documentação rigorosa de conformidade regulatória, análise de custos operacionais e avaliação de viabilidade logística para adoção em larga escala.

## Referências

- [ABRAMED 2024] ABRAMED (2024), O uso da IA na medicina diagnóstica brasileira – relatório de pesquisa, São Paulo: Abramed. Disponível em: <https://acrobat.adobe.com/id/urn:aaid:sc:VA6C2:5bdd30c9-0710-47df-b8e8-ee728facb828>.
- [Adler et al. 2023] Adler, J., Lenski, M., Tolios, A., Fares Taie, S., Sopic, M., Rajdl, D. and et al. (2023), "Digital competence in laboratory medicine", J Lab Med, v. 47, n. 4, p. 143-148. doi: [10.1515/labmed-2023-0021](https://doi.org/10.1515/labmed-2023-0021). [GS Search](#)
- [Aria and Cuccurullo 2017] Aria, M. and Cuccurullo, C. (2017), "bibliometrix: An R-tool for comprehensive science mapping analysis", J Informetr, v. 11, n. 4, p. 959-975. doi: [10.1016/j.joi.2017.08.007](https://doi.org/10.1016/j.joi.2017.08.007). [GS Search](#)
- [Baid, U. et al. 2024] Baid, U. and et al. (2024), "Pan-Cancer Tumor Infiltrating Lymphocyte Detection based on Federated Learning", In: 2024 IEEE International Conference on Big Data (BigData), Washington, DC, USA, p. 7640-7647. doi: [10.1109/BigData62323.2024.10825083](https://doi.org/10.1109/BigData62323.2024.10825083). [GS Search](#)
- [Cadamuro et al. 2024] Cadamuro, J., Carobene, A., Cabitza, F., Debeljak, Z., de Bruyne, S., van Doorn, W. and et al. (2024), "A comprehensive survey of artificial intelligence adoption in European laboratory medicine: Current utilization and prospects", Clin Chem Lab Med. doi: [10.1515/cclm-2024-1016](https://doi.org/10.1515/cclm-2024-1016). [GS Search](#)
- [Chen, B. et al. 2024] Chen, B., Zeng, H., Xiang, T., Guo, S., Zhang, T. and Liu, Y. (2024), "ESB-FL: Efficient and Secure Blockchain-Based Federated Learning With Fair Payment", IEEE Trans Big Data, v. 10, n. 6, p. 761-774. doi: [10.1109/TBDATA.2022.3177170](https://doi.org/10.1109/TBDATA.2022.3177170). [GS Search](#)

- [Düsing, C. et al. 2024] Düsing, C., Cimiano, P., Rehberg, S., Scherer, C., Kaup, O., Köster, C. and et al. (2024), "Integrating federated learning for improved counterfactual explanations in clinical decision support systems for sepsis therapy", *Artif Intell Med*, v. 157, p. 102982. doi: [10.1016/j.artmed.2024.102982](https://doi.org/10.1016/j.artmed.2024.102982). [GS Search](#)
- [Fang, C. et al. 2024] Fang, C. and et al. (2024), "Decentralised, collaborative, and privacy-preserving machine learning for multi-hospital data", *eBioMedicine*, v. 101, p. 105006. doi: [10.1016/j.ebiom.2024.105006](https://doi.org/10.1016/j.ebiom.2024.105006) . [GS Search](#)
- [Fantonelli et al. 2020] Fantonelli, R. L., Celuppi, I. C., Oliveira, F. M., Burigo, F., Dalmarco, E. M. and Wazlawick, R. S. (2020), "Lei geral de proteção de dados e a interoperabilidade na saúde pública", *J Health Inform*, Número Especial SBIS - Dezembro, p. 166-171. Disponvel em: [Lei geral de proteção de dados e a interoperabilidade na saúde pública](#) [General data protection law and interoperability in public health](#) [Ley general de protección de datos e interoperabilidad en salud pública](#).
- [Feng et al. 2025] Feng, Y., Guo, Y., Hou, Y., Wu, Y., Lao, M., Yu, T., Liu, G. (2025). A survey of security threats in federated learning. *Complex & Intelligent Systems*. 2025 Jan 29;11:165. doi: <https://doi.org/10.1007/s40747-024-01664-0>. [GS Search](#)
- [Fruchterman and Reingold 1991] Fruchterman, T. M. J. and Reingold, E. M. (1991), "Graph drawing by force-directed placement", *Softw Pract Exp*, v. 21, n. 11, p. 1129-1164. doi: <https://doi.org/10.1002/spe.4380211102> . [GS Search](#)
- [Guembe et al. 2024] Guembe, B., Misra, S., Azeta, A. (2024). Privacy Issues, Attacks, Countermeasures and Open Problems in Federated Learning: A Survey. *Applied Artificial Intelligence*. 2024;38(1):e2410504. <https://doi.org/10.1080/08839514.2024.2410504>. [GS Search](#)
- [Gulati, S. et al. 2024] Gulati, S., Guleria, K. and Goyal, N. (2024), "Collaborative, Privacy-Preserving Federated Learning Framework for the Detection of Diabetic Eye Diseases", *SN Comput Sci*, v. 5, p. 1100. doi: [10.1007/s42979-024-03462-4](https://doi.org/10.1007/s42979-024-03462-4). [GS Search](#)
- [Hajder, M. et al. 2021] Hajder, M., Hajder, P., Gil, T., Krzywda, M., Kolbusz, J. and Liput, M. (2021), "Architecture and organization of a Platform for diagnostics, therapy and post-covid complications using AI and mobile monitoring", *Procedia Comput Sci*, v. 192, p. 3711-3721. doi: [10.1016/j.procs.2021.09.145](https://doi.org/10.1016/j.procs.2021.09.145). [GS Search](#)
- [He, P. et al. 2023] He, P. and et al. (2023), "Low-Latency Federated Learning via Dynamic Model Partitioning for Healthcare IoT", *IEEE J Biomed Health Inform*, v. 27, n. 10, p. 4684-4695. doi: [10.1109/JBHI.2023.3298446](https://doi.org/10.1109/JBHI.2023.3298446). [GS Search](#)
- [Iturry et al. 2021] Díaz Iturry, M., Alves-Souza, S. N., Ito, M. and da Silva, S. A. (2021), "Data Quality in health records: A literature review", In: 2021 16th Iberian Conference on Information Systems and Technologies (CISTI), Chaves, Portugal, p. 1-6. doi: [10.23919/CISTI52073.2021.9476536](https://doi.org/10.23919/CISTI52073.2021.9476536). [GS Search](#)

- [Jamshidi, M. et al. 2025] Jamshidi, M., Hoang, D. T., Nguyen, D. N., Niyato, D. and Warkiani, M. E. (2025), "Revolutionizing biological digital twins: Integrating internet of bio-nano things, convolutional neural networks, and federated learning", *Comput Biol Med*, v. 189, p. 109970. doi: [10.1016/j.compbimed.2025.109970](https://doi.org/10.1016/j.compbimed.2025.109970). [GS Search](#)
- [Jimenez-Gutierrez et al. 2025] Jimenez-Gutierrez, D. M., Falkouskaya, Y., Hernandez-Ramos, J. L., Anagnostopoulos, A., Chatzigiannakis, I., Vitaletti, A. (2025). On the Security and Privacy of Federated Learning: A Survey with Attacks, Defenses, Frameworks, Applications, and Future Directions. arXiv:2508.13730. 2025 Aug 19. <https://arxiv.org/abs/2508.13730>. [GS Search](#)
- [Kitchenham et al. 2007] Kitchenham, B., Charters, S., Budgen, D., Brereton, P., Turner, M., Linkman, S., Jørgensen, M., Mendes, E., Visaggio, G. (2007), Guidelines for performing systematic literature reviews in software engineering, EBSE Technical Report. Software Engineering Group –Keele University, UK. Disponível em: [slr.key](https://www.slr.key.org) .
- [Kodali et al. 2024] Kodali, R. K., Mittadoddi, V. P., Ranga, H. and Boppana, L. (2024), "Machine Learning in Laboratory Diagnosis", In: TENCON 2024 - 2024 IEEE Region 10 Conference (TENCON), Singapore, p. 1445-1449. doi: [10.1109/TENCON61640.2024.10902716](https://doi.org/10.1109/TENCON61640.2024.10902716). [GS Search](#)
- [Kumar, D. et al. 2025] Kumar, D., Verma, C. and Illés, Z. (2025), "Federated learning with explainable AI for liver disease prediction: A privacy-preserving approach", *Intell Based Med*, v. 12, p. 100285. doi: [10.1016/j.ibmed.2025.100285](https://doi.org/10.1016/j.ibmed.2025.100285). [GS Search](#)
- [Li, Z. et al. 2024] Li, Z., Li, M., Li, A. and Lin, Z. (2024), "Blockchain-based collaborative data analysis framework for distributed medical knowledge extraction", *Comput Ind Eng*, v. 190, p. 110099. doi: [10.1016/j.cie.2024.110099](https://doi.org/10.1016/j.cie.2024.110099). [GS Search](#)
- [Lippi and Plebani 2025] Lippi, G. and Plebani, M. (2025), "Lights and shadows of artificial intelligence in laboratory medicine", *Adv Lab Med*, v. 6, n. 1, p. 1-3. doi: [10.1515/almed-2025-0024](https://doi.org/10.1515/almed-2025-0024). [GS Search](#)
- [Liu, X. et al. 2025] Liu, X., Li, S., Zhu, Q. and et al. (2025), "Interpretable Semi-federated Learning for Multimodal Cardiac Imaging and Risk Stratification: A Privacy-Preserving Framework", *J Digit Imaging Inform Med*. doi: [10.1007/s10278-025-01643-y](https://doi.org/10.1007/s10278-025-01643-y). [GS Search](#)
- [Lu, M. Y. et al. 2022] Lu, M. Y., Chen, R. J., Kong, D., Lipkova, J., Singh, R., Williamson, D. F. K. and et al. (2022), "Federated learning for computational pathology on gigapixel whole slide images", *Med Image Anal*, v. 76, p. 102298. doi: [10.1016/j.media.2021.102298](https://doi.org/10.1016/j.media.2021.102298). [GS Search](#)
- [Ma, T. et al. 2025] Ma, T., Luo, X., Tan, R. and Gao, H. (2025), "Privacy and Fairness-Guaranteed Federated Preference Optimization for Large Language Models in Internet of Medical Things", *IEEE Trans Consum Electron*. doi: [10.1109/TCE.2025.3595092](https://doi.org/10.1109/TCE.2025.3595092). [GS Search](#)



- [Master et al. 2023] Master, S. R., Badrick, T. C., Bietenbeck, A. and Haymond, S. (2023), "Machine Learning in Laboratory Medicine: Recommendations of the IFCC Working Group", Clin Chem, v. 69, n. 7, p. 690-698. doi: [10.1093/clinchem/hvad055](https://doi.org/10.1093/clinchem/hvad055). [GS Search](#)
- [McMahan et al. 2017] McMahan, B., Moore, E., Ramage, D., Hampson, S. and Arcas, B. A. y. (2017), "Communication-efficient learning of deep networks from decentralized data", In: Proceedings of the 20th International Conference on Artificial Intelligence and Statistics (AISTATS), Fort Lauderdale, FL, USA, p. 1273-1282. [GS Search](#)
- [Mnasri, S. et al. 2025] Mnasri, S., Salah, D. and Idoudi, H. (2025), "A hybrid blockchain and federated learning attention-based BERT transformer framework for medical records management", J Supercomput, v. 81, p. 317. doi: [10.1007/s11227-024-06816-0](https://doi.org/10.1007/s11227-024-06816-0). [GS Search](#)
- [Page et al. 2021] Page, M. J., McKenzie, J. E., Bossuyt, P. M., Boutron, I., Hoffmann, T. C., Mulrow, C. D. and et al. (2021), "The PRISMA 2020 statement: an updated guideline for reporting systematic reviews", BMJ, v. 372, p. n71. doi: [10.1136/bmj.n71](https://doi.org/10.1136/bmj.n71). [GS Search](#)
- [Park, E. et al. 2022] Park, E., Nam, H. S. and Song, J. W. (2022), "Federated Learning Models using Flow Cytometry Data of Blood Test in Medical Decision Support", In: 2022 IEEE International Conference on Big Data (Big Data), Osaka, Japan, p. 6793-6795. doi: [10.1109/BigData55660.2022.10020379](https://doi.org/10.1109/BigData55660.2022.10020379). [GS Search](#)
- [Perales, E. et al. 2025] Perales, E., Verdy-Ricard, R., Labiod, M. A., Bendiab, G. and Chenoune, Y. (2025), "Federated Learning for Securing Medical Imaging Against Deepfakes in 6G Smart Hospitals", In: 2025 IEEE International Conference on Cyber Security and Resilience (CSR), Chania, Crete, Greece, p. 1100-1105. doi: [10.1109/CSR64739.2025.11130027](https://doi.org/10.1109/CSR64739.2025.11130027). [GS Search](#)
- [Prabh et al. 2024] Prabha, B. V., Manikanda Kumaran, K., Manikandan, S. and Murugan, S. P. (2024), "A Comparative Analysis of Machine Learning Algorithms for Healthcare Applications", In: 2024 4th International Conference on Advancement in Electronics & Communication Engineering (AECE), GHAZIABAD, India, p. 214-218. doi: [10.1109/AECE62803.2024.10911687](https://doi.org/10.1109/AECE62803.2024.10911687). [GS Search](#)
- [Qiu, W. et al. 2024] Qiu, W. and et al. (2024), "Heart Sound Abnormality Detection From Multi-Institutional Collaboration: Introducing a Federated Learning Framework", IEEE Trans Biomed Eng, v. 71, n. 10, p. 2802-2813. doi: [10.1109/TBME.2024.3393557](https://doi.org/10.1109/TBME.2024.3393557). [GS Search](#)
- [Sai, S. et al. 2024] Sai, S., Hassija, V., Chamola, V. and Guizani, M. (2024), "Federated Learning and NFT-Based Privacy-Preserving Medical-Data-Sharing Scheme for Intelligent Diagnosis in Smart Healthcare", IEEE Internet Things J, v. 11, n. 4, p. 5568-5577. doi: [10.1109/JIOT.2023.3308991](https://doi.org/10.1109/JIOT.2023.3308991). [GS Search](#)

- [Shahnazeer, C. K. and Sureshkumar, G. 2025] Shahnazeer, C. K. and Sureshkumar, G. (2025), "Federated Transfer Learning Framework for Multi-disease Prediction", *Procedia Comput Sci*, v. 258, p. 830-838. doi: [10.1016/j.procs.2025.04.315](https://doi.org/10.1016/j.procs.2025.04.315). [GS Search](#)
- [Sheller et al. 2020] Sheller, M. J., Edwards, B., Reina, G. A. and et al. (2020), "Federated learning in medicine: facilitating multi-institutional collaborations without sharing patient data", *Sci Rep*, v. 10, p. 12598. doi: [10.1038/s41598-020-69250-1](https://doi.org/10.1038/s41598-020-69250-1). [GS Search](#)
- [Wei, M. et al. 2025] Wei, M., Yu, W. and Chen, D. (2025), "AccDFL: Accelerated Decentralized Federated Learning for Healthcare IoT Networks", *IEEE Internet Things J*, v. 12, n. 5, p. 5329-5345. doi: [10.1109/JIOT.2024.3486122](https://doi.org/10.1109/JIOT.2024.3486122). [GS Search](#)
- [Xiao 2023] Xiao, F. (2023), "Application of Computer Big Data and Cloud Computing Precision Inspection in Laboratory Treatment", In: 2023 International Conference on Telecommunications, Electronics and Informatics (ICTEI), Lisbon, Portugal, p. 779-782. doi: [10.1109/ICTEI60496.2023.00017](https://doi.org/10.1109/ICTEI60496.2023.00017). [GS Search](#)
- [Zhang, H. et al. 2025] Zhang, H., Chen, M., Liu, Y., Luo, G. and Zhu, Y. (2025), "Non-IID Medical Image Segmentation Based on Cascaded Diffusion Model for Diverse Multi-Center Scenarios", *IEEE J Biomed Health Inform*, v. 29, n. 7, p. 5042-5055. doi: [10.1109/JBHI.2025.3549029](https://doi.org/10.1109/JBHI.2025.3549029). [GS Search](#)
- [Zhang, L. et al. 2023] Zhang, L., Xu, J., Vijayakumar, P., Sharma, P. K. and Ghosh, U. (2023), "Homomorphic Encryption-Based Privacy-Preserving Federated Learning in IoT-Enabled Healthcare System", *IEEE Trans Netw Sci Eng*, v. 10, n. 5, p. 2864-2880. doi: [10.1109/TNSE.2022.3185327](https://doi.org/10.1109/TNSE.2022.3185327). [GS Search](#)
- [Zhao et al. 2025] Zhao, J.C. and et al. (2025). The Federation Strikes Back: A Survey of Federated Learning Privacy Attacks, Defenses, Applications, and Policy Landscape. *ACM Computing Surveys*. 2025;57(9):230. <https://doi.org/10.1145/3724113> . [GS Search](#)
- [Zou, W. et al. 2025] Zou, W., Zhang, R. and Xun, Y. (2025), "On a Federated-Learning-Based Computation Offloading Strategy for Nonterrestrial-Network-Assisted Internet of Medical Things", *IEEE Internet Things J*, v. 12, n. 12, p. 21280-21289. doi: [10.1109/JIOT.2025.3546812](https://doi.org/10.1109/JIOT.2025.3546812). [GS Search](#)