

Número 116 | Julho 2026



Radar

A revista de cibersegurança

A identidade como novo perímetro: a confiança começa em cada identidade

Por Hans Vigil Navas

Durante anos, a cibersegurança se apoiou em uma ideia simples: proteger o perímetro. Havia uma rede interna, usuários conhecidos, sistemas sob controle e uma fronteira relativamente clara entre o que era confiável e o que vinha de fora. Firewalls, segmentação de rede e controles de entrada pareciam suficientes para defender aquilo que a organização considerava crítico. Mas essa fronteira deixou de ser evidente.

Cloud, trabalho híbrido, APIs, terceiros, automação industrial e inteligência artificial transformaram o acesso em uma conversa permanente entre pessoas, aplicações, máquinas e sistemas que operam dentro e fora da organização. **Hoje, o risco já não surge apenas quando alguém tenta atravessar uma barreira.** Muitas vezes, o risco já está dentro do ambiente, usando uma credencial válida, um token ativo, um certificado esquecido ou uma conta de serviço com mais privilégios do que o necessário.

Nesse novo cenário, a identidade se tornou o novo perímetro. Isso já não é mais uma mera tendência. É uma realidade operacional. **Cada identidade representa uma possibilidade de acesso, uma decisão de confiança e, ao mesmo tempo, uma porta potencial para os ativos críticos da organização.** Algumas identidades pertencem a pessoas. Outras, a aplicações, serviços, dispositivos, processos automatizados ou agentes capazes de executar tarefas em nome de outra entidade. Muitas dessas identidades não aparecem nos inventários tradicionais nem são revisadas com o mesmo rigor aplicado a um usuário humano.

O desafio é que a confiança digital se tornou mais difícil de observar. Uma organização pode contar com boas ferramentas de monitoramento, criptografia, resposta a incidentes e proteção de rede. Mas, se não sabe quem acessa, o que acessa, com quais privilégios, de onde, por quanto tempo e em nome de quem atua, sua postura de segurança continua incompleta. **A identidade deixa de ser um componente administrativo de TI para se tornar uma capacidade crítica de governança, resiliência e continuidade.**

No setor bancário, uma identidade comprometida pode se transformar em fraude. No setor de seguros, em exposição de informações sensíveis. Na mineração, pode causar interrupção da produção. Em setores industriais, em impacto operacional. Em todos os casos, a origem costuma ser a mesma: uma confiança mal governada.

Para os setores financeiro e segurador, essa realidade é especialmente relevante. A identidade está no centro da prevenção a fraudes, da proteção dos dados de clientes, do controle de acessos privilegiados, da segurança de canais digitais, do uso de APIs e da relação com terceiros. Nesses ambientes, não basta implementar controles. É preciso demonstrar que funcionam, que são revisados, que geram evidências e que permitem responder a auditores, reguladores e conselhos de administração.

Na mineração e em setores industriais, o impacto pode ser ainda mais tangível. Uma identidade comprometida nem sempre termina em vazamento de dados. Também pode abrir caminho para acessos remotos não autorizados, contas compartilhadas, credenciais persistentes de fornecedores, estações de engenharia ou sistemas que sustentam processos críticos. Nesses setores, a identidade não protege apenas a informação. Esse controle também protege disponibilidade, continuidade operacional, segurança física e capacidade produtiva.

Por isso, falar de identidade como novo perímetro não significa falar apenas de uma nova ferramenta de IAM. Significa repensar a forma como a organização entende a confiança. Cada acesso deve ser explícito, verificável, mínimo, rastreável e revogável. Cada identidade deve ter um propósito, um responsável, um ciclo de vida e controles associados. Cada privilégio deve responder a uma necessidade real do negócio. E cada exceção deve ser tratada como um risco, não como um hábito operacional.

Os principais frameworks internacionais vêm apontando esse caminho sob diferentes perspectivas. O NIST aborda o tema a partir de Zero Trust, sem confiança implícita e com verificação contínua. A ISO/IEC 27001 traduz essa lógica em gestão de riscos, controles de acesso, relação com terceiros e melhoria contínua. A OWASP evidencia o problema a partir de falhas recorrentes em autenticação, autorização, segurança de APIs e identidades não humanas. **Embora partam de abordagens diferentes, todas apontam para a mesma conclusão: a confiança deve ser governada.**

A gestão de identidades não humanas se torna, então, um ponto crítico. Em muitas organizações, essas identidades já superam amplamente as identidades humanas em volume e criticidade. Certificados digitais, segredos, tokens, contas de automação, integrações entre plataformas, serviços cloud, dispositivos conectados e processos automatizados operam em segundo plano, sustentando funções essenciais do negócio. Quando uma dessas identidades é comprometida, o atacante não precisa romper o perímetro. Basta agir como se fizesse parte dele.

A esse cenário soma-se uma nova camada de complexidade, os agentes de identidade, ou **ID Agents**. **À medida que a inteligência artificial evolui para modelos mais autônomos, surgem agentes capazes de atuar em nome de uma pessoa, uma aplicação ou um processo.** Esses agentes podem consumir APIs, executar tarefas, tomar decisões assistidas e operar com certo grau de independência. A pergunta, portanto, já não é apenas “quem acessa?”, mas “em nome de quem esse agente atua, com qual autoridade, sob quais limites e com qual rastreabilidade?”.

A autonomia sem governança pode transformar a identidade em um ponto cego. Com os controles adequados, a autonomia pode fortalecer a eficiência sem comprometer a segurança. Esse é o desafio dos CISOs, que precisam integrar governança de acessos, gestão de privilégios, proteção de segredos, segurança de APIs, controle de terceiros e supervisão contínua sob uma mesma visão de risco. Em bancos e seguradoras, isso se traduz em prevenção a fraudes, conformidade, evidências e confiança do cliente. Na mineração e em setores industriais, em continuidade, disponibilidade e proteção de ambientes físico-digitais. Em ambos os casos, a identidade deixa de ser um controle isolado para se tornar um eixo transversal de defesa e resiliência.

Nesta edição da Radar, abordamos justamente essa mudança de paradigma. A segurança já não pode depender de fronteiras que o negócio deixou para trás. Deve ser construída em torno de identidades verificáveis, humanas e não humanas, gerenciadas ao longo de todo o seu ciclo de vida e alinhadas a padrões, riscos e regulamentações.

O novo perímetro não está na rede, mas em cada identidade com capacidade de agir. Algumas são humanas, outras são máquinas e outras serão agentes autônomos. Todas exigem governança.

Porque, na era da hiperconexão, a confiança já não se concede. A confiança se verifica, se limita, se monitora e se revoga.



Hans Vigil Navas
Cybersecurity Manager

As chaves de acesso: a identidade como principal vetor de ataque

Cibercrônica por Nelson Andrés Barboza Landinez

Durante décadas, a segurança corporativa foi construída sobre uma premissa simples: havia um ambiente interno e um ambiente externo. O firewall funcionava como barreira, a VPN como ponte de acesso e os antivírus como controles distribuídos pelos endpoints. Era um modelo previsível, desenhado para um mundo em que as ameaças vinham de fora e os ativos confiáveis permaneciam dentro da organização. Mas esse mundo desapareceu. O trabalho remoto abriu novos pontos de entrada. As integrações SaaS ampliaram a dependência de fornecedores externos. Hoje, não existe apenas um “dentro”. Existem identidades.

De acordo com o relatório de resposta a incidentes da Unit 42 para 2026, 65% dos acessos iniciais em ciberataques ocorreram por meio de técnicas baseadas em identidade, permitindo acesso não autorizado, movimentação lateral e escalonamento de privilégios. Técnicas como phishing, credenciais roubadas, ataques de força bruta, configurações incorretas de IAM e atividade de insiders estão no centro das violações modernas. Em um cenário no qual organizações investem volumes significativos de dinheiro e recursos em ciberdefesa, os atacantes escolhem o caminho mais direto e entram pela porta da frente. O DBIR 2025 da Verizon confirma essa realidade. Segundo o relatório, 22% das violações começam com credenciais roubadas, um aumento de 34% em relação ao ano anterior. Não estão arrombando sistemas. Estão fazendo login.

A rápida adoção de infraestrutura cloud, SaaS, contas usadas por aplicações, serviços e processos automatizados criou um ambiente digital no qual cada integração representa uma possível rota de acesso. O mesmo relatório da Unit 42 demonstrou que, em uma amostra de 680.000 identidades em cloud, 99% dos usuários, roles e serviços tinham permissões excessivas, incluindo acessos que não haviam sido utilizados por mais de 60 dias.

O atacante não precisa escalar pela força. Basta avançar por acessos que ninguém revogou. Soma-se a isso a ameaça crescente da computação quântica aos esquemas atuais de criptografia e a redução prevista do ciclo de vida dos certificados TLS. O tempo está correndo e a superfície de exposição continua crescendo.

O relatório da CyberArk de 2025 ouviu 1.200 líderes de segurança em organizações dos Estados Unidos, Reino Unido, Austrália, França e Alemanha, revelando que a maioria enfrenta dificuldades sérias para gerenciar de forma eficaz as identidades de máquina.

Os casos de 2025 ilustram essa realidade com contundência. A Marks & Spencer, varejista multinacional britânica, foi alvo de um ciberataque atribuído ao grupo “Scattered Spider” em abril de 2025. A empresa confirmou a criptografia de múltiplos servidores, com impacto significativo em sua infraestrutura.

A violação de segurança não ocorreu de forma isolada. Foi resultado de um processo no qual a gestão de identidades se mostrou o elo mais fraco. Durante a intrusão, os atores maliciosos obtiveram acesso a um arquivo NTDS.dit, o que permitiu extrair hashes de senhas e facilitou o escalonamento de privilégios na rede. Embora o mecanismo exato do comprometimento inicial não tenha sido confirmado, diversas fontes indicam que o grupo Scattered Spider costuma usar campanhas de phishing, vishing e tentativas de fadiga de MFA para obter acesso a contas corporativas. Posteriormente, a implantação do criptografador do ransomware DragonForce afetou a continuidade de múltiplos serviços corporativos e comerciais, além de viabilizar acesso não autorizado a informações de clientes, conforme confirmado pela organização.

Meses depois, entre 8 e 18 de agosto, outra campanha de roubo de informações teve como alvo a plataforma de automação de compras Salesloft, especialmente instâncias de clientes Salesforce, por meio do comprometimento de tokens OAuth associados ao agente de inteligência artificial Drift.

Segundo a Cloudflare, a rota de ataque envolveu enumeração de objetos, reconhecimento do escopo do ambiente comprometido e download de informações por interfaces legítimas do Salesforce.

Por fim, como parte de uma operação cuidadosamente executada, houve a exclusão de jobs para ocultar as ações realizadas. O agente de ameaça, identificado pelo Google Threat Intelligence Group (GTIG) como UNC6395, teria afetado cerca de 700 organizações, incluindo instituições financeiras, provedores de saúde, empresas de tecnologia e agências governamentais.

Cory Michal, CSO da AppOmni, afirmou que muitas das organizações comprometidas são empresas de software e segurança, o que indica que a campanha pode representar um movimento inicial para futuros ataques à cadeia de suprimentos.

Não foi uma falha técnica em endpoints nem na rede. Foi uma falha na identidade do agente. Tokens OAuth de longa duração, com escopos permissivos, pertencentes a um agente confiável que ninguém monitorava. As organizações confiam nas integrações como confiam nos colaboradores, mas não aplicam governança a essas integrações. Agentes são identidades. E identidades sem governança são portas abertas.

Apenas algumas semanas depois, a Jaguar Land Rover revelou ter sofrido um ciberataque devastador em 31 de agosto de 2025, que custou à companhia 196 milhões de libras. Além disso, de acordo com o The Cyber Monitoring Center, o ataque gerou um impacto estimado em 1,9 bilhão de libras na economia do Reino Unido, devido à paralisação total das operações globais de manufatura por cinco semanas. O ataque, atribuído ao grupo Scattered LAPSUS\$ Hunters, cuja filosofia é “log in, not hack in”, foi possível graças ao roubo de credenciais por meio de um infostealer, um tipo de malware projetado para extrair senhas, cookies de sessão e outras informações sensíveis de equipamentos comprometidos. Neste caso, as credenciais pertenciam a um terceiro com acesso ao sistema Jira da JLR.

Em 1º de setembro, a equipe de TI da JLR detectou uma intrusão na rede e tomou uma medida drástica ao desligar todos os sistemas para conter o dano. Para comprovar o acesso, o grupo de hackers publicou imagens dos sistemas internos no Telegram, expondo acesso a domínios como jlrint.com. No entanto, análises posteriores revelaram que a intrusão não foi resultado de uma vulnerabilidade zero-day. O ataque combinou táticas como engenharia social, abuso de credenciais, detecção deficiente e segmentação fraca entre a rede corporativa e os sistemas de produção.

Casos como esses não ocorrem no vácuo. A inteligência artificial desempenha um papel determinante na velocidade com que essas campanhas de roubo de identidade são executadas. A análise de 32 milhões de e-mails de phishing documentada no Darktrace Annual Threat Report 2026 mostrou uma tendência clara. Os ataques de phishing estão cada vez mais sofisticados e potencialmente mais eficazes quando esses e-mails chegam a usuários com privilégios elevados.

Essa conclusão não parece exagerada, considerando que 70% desses e-mails passaram pelos controles de autenticação DMARC, permitindo driblar controles de monitoramento e detecção. Nathaniel Jones, vice-presidente de segurança e estratégia de IA da Darktrace, afirma que “o perímetro tradicional de defesa foi criado para um mundo em que os atacantes precisavam invadir”.

Embora os agentes mal-intencionados tenham historicamente baseado suas estratégias em exploits técnicos, o desenvolvimento da IA permitiu operar com maior precisão, adotando técnicas como vishing e deepfakes, com foco nas identidades como vetor de acesso inicial, mais do que na infraestrutura em si.

Mitigações efetivas, como a implantação de MFA resistente a phishing, a governança de identidades, a implementação e a aplicação efetiva do princípio do menor privilégio, além de soluções de monitoramento integradas à inteligência artificial para detectar comportamentos anômalos, como tokens criados pela primeira vez a partir de novas regiões ou padrões incomuns no consumo de APIs, são elementos fundamentais de uma estratégia moderna de defesa.

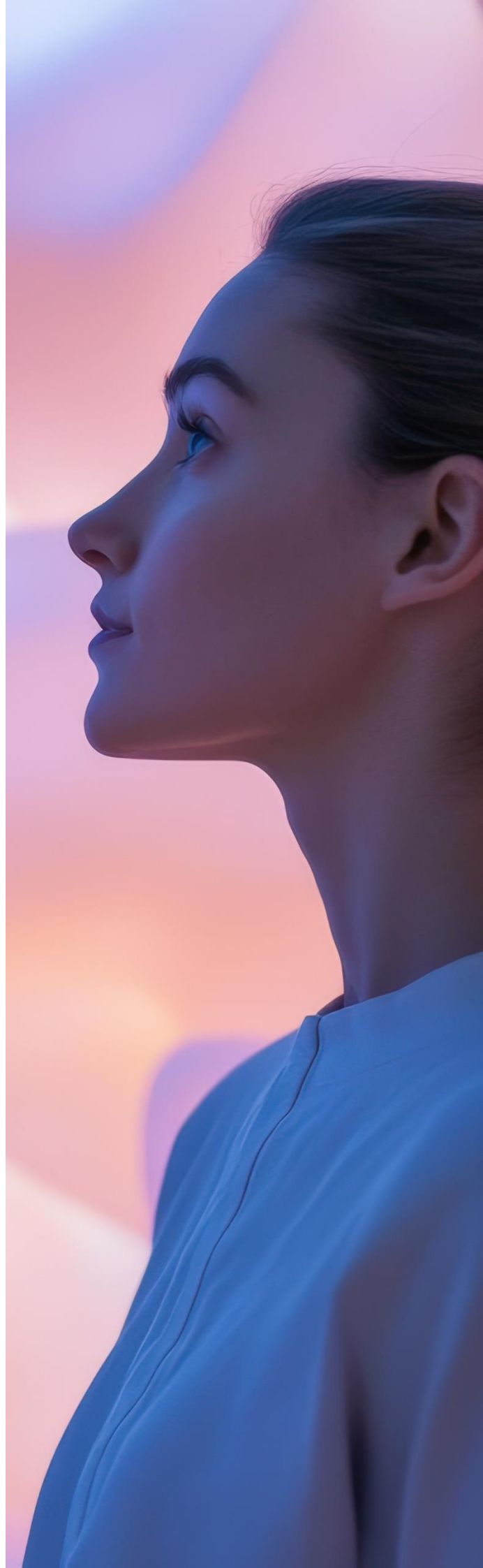
Também é indispensável contar com resposta a incidentes baseada em playbooks que considerem a possibilidade de intrusos já estarem dentro do sistema. Da mesma forma, controlar a reutilização de senhas e identificar identidades expostas que já circulam online são práticas essenciais neste novo cenário, em que as identidades passaram a ocupar uma posição central no desenvolvimento tecnológico.

Embora as senhas nunca tenham sido seguras, neste momento os fatos falam por si. Nesta nova era transformada pela inteligência artificial, as senhas como único mecanismo de autenticação se mostraram insuficientes diante das ameaças atuais. Chegou o momento de migrar para mecanismos de autenticação mais robustos. A confidencialidade, a integridade e a disponibilidade dos sistemas não podem depender apenas de barreiras projetadas para resistir a ataques externos. Também é preciso proteger os principais pontos de acesso e monitorar quem tem entrada no ambiente interno.

Ao longo dos anos, milhões foram investidos para reforçar defesas digitais, sob a convicção de que o perigo viria de fora. No entanto, alguns dos incidentes mais catastróficos dos últimos meses mostram uma realidade diferente. Os atacantes já não precisam derrubar portas. Basta obter as chaves. Na era da cloud, dos agentes inteligentes e da inteligência artificial, cada usuário, serviço, token, aplicação e máquina possui uma identidade própria. Essa identidade se tornou o novo perímetro de segurança.



Nelson Andrés Barboza Landinez
Cybersecurity Engineer



A identidade: o risco que evoluiu junto com a IA

Artículo por Jaime Andrés Tovar Prieto

Quantos agentes de IA estão operando agora mesmo na sua organização com acesso amplo e sem identidade verificável? É possível saber qual identidade esses agentes estão usando? Há quanto tempo operam com essa identidade atribuída? Nos últimos anos, a inteligência artificial consolidou sua presença no ecossistema tecnológico corporativo e evoluiu para além dos modelos de IA Generativa ou conversacionais padronizados, projetados para interpretar um contexto e gerar respostas sob supervisão constante de um ser humano. Esse modelo é conhecido como Human-in-the-Loop (HITL). A etapa seguinte dessa evolução são os agentes de IA, componentes de software autônomos que não apenas interpretam entradas e planejam estratégias, mas também executam tarefas tangíveis por meio da comunicação com interfaces de programação de aplicações (Application Programming Interfaces, APIs). Esses agentes interagem com bancos de dados, outros agentes e subagentes em fluxos de trabalho coordenados, com intervenção humana mínima ou inexistente em cada etapa individual do processo.

A governança convencional da IA se concentrou em mitigar o risco associado às respostas do modelo, avaliando se o resultado é preciso, justo, ético ou livre de vieses e alucinações, sob a premissa de que um ser humano revisará os resultados antes de tomar uma decisão ou executar uma ação. No entanto, a governança de Agentic AI exige uma mudança estrutural e conceitual. O controle de risco deve se orientar para a ação. Um agente dotado de autonomia pode iniciar transações financeiras complexas, atualizar registros de clientes em sistemas corporativos ou aplicar configurações em dispositivos de segurança de infraestrutura crítica sem esperar nenhuma confirmação humana interativa. Portanto, a ausência de controles rigorosos e dinâmicos em tempo de execução pode resultar em escalonamentos de privilégios não intencionais, vazamento massivo de dados e alterações sistêmicas capazes de comprometer a continuidade do negócio.

A nova fronteira do risco: da resposta à ação

Esses cenários geraram alertas de alto nível no setor. A Gartner projeta que, até 2027, 40% das empresas serão obrigadas a degradar ou desmantelar completamente suas implantações de agentes autônomos devido a fragilidades nos modelos de governança. Organizações que não conseguem distinguir entre a capacidade técnica de um agente para executar uma ação e o alcance legítimo das permissões concedidas enfrentam um risco operacional iminente e insustentável.

Diante desse cenário, é fundamental analisar a nova fronteira do risco, na qual a identidade emerge como o componente crítico – em evolução – no contexto da IA.

A Okta aponta, no artigo “What is AI Agent Identity? Securing Autonomous Systems”, que, em média, uma organização gerencia 45 identidades de máquina, ou identidades não humanas (Non-Human Identities, NHIs), para cada usuário humano.

No entanto, enquanto as NHIs tradicionais operam com parâmetros previsíveis e sequências de comandos predefinidas, sem consciência do ambiente, os agentes de IA têm capacidade de raciocinar, adaptar suas estratégias de forma iterativa conforme o contexto, reter informações para aprimorar a tomada de decisões e executar ações autônomas com impacto real nos ambientes. Aplicar os modelos rígidos de identidade atualmente vigentes à Agentic AI pode gerar brechas de segurança de alto impacto.

Gestão de identidades e acessos padrão vs. identidade de agentes: diferenças fundamentais

As diferenças estruturais entre a gestão convencional de identidade e os requisitos específicos da identidade de agentes se manifestam em três dimensões críticas, que alteram fundamentalmente o ciclo de vida da identidade:

- **Vigência:** A identidade de um usuário humano pode permanecer ativa por dias, meses ou anos e exige validações interativas, como biometria, senhas, tokens ou mecanismos de duplo fator. Os agentes, por sua vez, operam de forma contínua e sem intervenção humana constante. Suas identidades podem ser efêmeras e existir apenas durante os segundos necessários para concluir uma transação.
- **Autorização delegada:** Contas humanas atuam em nome próprio. Já os agentes de IA atuam em nome de um humano, sistema ou serviço, em um modelo em nome de outro (On-Behalf-Of, OBO). Isso exige cadeias de delegação criptograficamente rigorosas e trocas de tokens capazes de preservar a responsabilização e a rastreabilidade forense, mantendo cada ação vinculada à identidade que delegou a autorização.

- **Controle de acesso:** A atribuição estática de perfis de acesso não consegue acompanhar o ritmo de um agente que encadeia tarefas complexas e toma milhares de decisões de acesso por minuto. Os agentes exigem avaliação contínua do contexto, do risco ambiental e dos atributos específicos de cada tarefa, instrumentada por meio de Policy-as-Code, ou políticas como código.

Vetores de risco em ambientes de IA não governados

Quando a velocidade de implementação supera a maturidade da governança de IA, cenário altamente provável no contexto atual, as organizações ficam expostas a riscos estruturais derivados de NHIs não gerenciadas.

Os principais vetores são:

- **Acesso persistente não governado:** Agentes que operam com credenciais permanentes ou tokens de acesso de longa duração se tornam alvos de alto valor. Por exemplo, quando um modelo é comprometido por meio de uma injeção de prompts maliciosa, o atacante herda todo o acesso permanente do agente e pode operar livremente dentro do perímetro corporativo.
- **Escalonamento de privilégios e deriva de autoridade (Privilege Drift):** Em ambientes dinâmicos, um agente pode invocar ferramentas ou assumir perfis de acesso que acumulam permissões nunca aprovadas explicitamente por um administrador de segurança. Essa deriva costuma se manifestar como uma expansão progressiva de autoridade ao longo do tempo, não como uma mudança abrupta e imediata nos resultados gerados.
- **Falta de atribuição e evasão de auditoria:** Quando um agente modifica um banco de dados usando uma conta de serviço genérica ou compartilhada, a ação perde seu contexto de delegação. A impossibilidade de distinguir algoritmicamente uma ação humana de uma ação sintética prejudica a responsabilização e viola normas rigorosas de conformidade.

Framework de governança para identidades de agentes

O primeiro passo para uma governança efetiva é estabelecer um framework de mensuração de risco aplicado aos agentes.

É fundamental evitar a implementação de um modelo de segurança uniforme para todos os agentes. Esse tipo de abordagem pode gerar tanto uma sobreproteção operacional, que limita a produtividade, quanto uma exposição ao risco decorrente da subestimação das capacidades reais do agente. As políticas de segurança devem ser proporcionais ao contexto e às capacidades específicas de cada entidade baseada em agentes.

A gestão de identidades e acessos (Identity and Access Management, IAM) tradicional foi projetada para gerenciar identidades humanas com ciclos de vida previsíveis e autenticação interativa. Em ambientes de Agentic AI, esse modelo é estruturalmente insuficiente. Os agentes não possuem credenciais próprias no sentido convencional, atuam por delegação e operam em velocidades que superam qualquer processo de revisão humana. Para superar essa lacuna, a arquitetura exige uma camada superior de abstração conhecida como Identity Fabric. Plataformas empresariais avançadas atuam como Identity Fabric, conectando diversas aplicações, provedores de identidade (Identity Providers, IDPs) e ambientes híbridos e multicloud sem exigir a reescrita das aplicações existentes.

A orquestração oferecida pelo Identity Fabric não se limita a conectar sistemas. Essa camada centraliza a gestão de identidade em seis dimensões operacionais:

- **Autenticação:** Fornecer serviços de login modernos e adaptativos, superando protocolos legados com vulnerabilidades conhecidas.
- **Controle de acesso (Access Control):** Impor políticas granulares e consistentes no nível de indivíduo, grupo ou máquina por meio de um plano de controle unificado.
- **Autorização:** Estabelecer regras precisas que determinem quais funções específicas um agente pode executar e sob quais condições exatas do ambiente.
- **Atributos:** Fornecer dados dinâmicos do usuário a partir de qualquer IDP em tempo de execução para apoiar decisões de segurança contextualizadas e personalizadas.
- **Administração:** Viabilizar a gestão centralizada e visual de políticas em todo o ecossistema.
- **Auditoria:** Garantir observabilidade profunda da atividade real em relação às permissões concedidas, demonstrando conformidade normativa de forma contínua.

Rumo a um framework de autenticação para agentes: SPIFFE

Os mecanismos habituais de autenticação em ambientes web, como tokens JSON Web Token (JWT), OAuth, credenciais estáticas e chaves de API, são insuficientes para gerenciar a autenticação em ecossistemas de Agentic AI. Para mitigar essa lacuna, surge um componente arquitetônico essencial, o **Secure Production Identity Framework for Everyone (SPIFFE)**.

O **SPIFFE** aborda especificamente um risco de segurança em camadas, frequentemente descrito como efeito de “boneca russa”. Quando um atacante compromete um subagente em um nível profundo da cadeia por meio de injeção de prompts, esse agente malicioso pode herdar os amplos privilégios concedidos ao agente principal, causando danos catastróficos por meio de uma técnica documentada como “contrabando de sessão agêntica” (Agent Session Smuggling).

O **SPIFFE** fornece um mecanismo padronizado para emitir identidades criptograficamente verificáveis ao software, com base em sua procedência e em seu contexto de execução, sem exigir intervenção humana para o provisionamento de credenciais. O elemento central é o **SPIFFE Verifiable Identity Document (SVID)**, uma credencial efêmera assinada criptograficamente como uma asserção.

Os **SVIDs** são administrados por meio do SPIFFE Runtime Environment (SPIRE), que garante sua rotação automática ao longo de todo o ciclo de vida do agente. Esse mecanismo elimina completamente o vetor de persistência maliciosa e materializa uma abordagem pura de Zero Trust.

Conclusão

A identidade em ambientes de Agentic AI não é um problema técnico menor. É o eixo central sobre o qual se constrói, ou se destrói, a confiança operacional de uma organização. Enquanto os agentes continuarem proliferando sem identidades verificáveis, sem ciclos de vida governados e sem rastreabilidade criptográfica, cada automação autônoma representará um vetor de risco aberto. O caminho para uma postura de segurança madura em IA exige adotar os princípios de Zero Trust não como diretrizes teóricas estáticas, mas como uma arquitetura executável, com identidades efêmeras, provisionadas Just-in-Time, com privilégios mínimos e auditabilidade contínua.



Jaime Andrés Tovar Prieto
Cybersecurity Architect

O risco invisível: como as identidades não humanas estão redefinindo a cibersegurança

Artigo por César Vega Calderón

As identidades não humanas constituem hoje um dos vetores de risco que mais crescem. A gestão adequada dessas identidades será determinante para a resiliência cibernética das organizações na era da cloud, da automação, da transformação digital e da inteligência artificial.

A adoção acelerada de serviços em cloud, automação, DevOps e inteligência artificial impulsionou o crescimento exponencial das identidades não humanas (Non-Human Identities, NHIs). Utilizadas por aplicações, interfaces de programação de aplicações (Application Programming Interfaces, APIs), contas de serviço, containers, dispositivos de Internet das Coisas (Internet of Things, IoT) e agentes de IA, essas identidades são essenciais para a operação moderna das organizações. Ao mesmo tempo, representam uma das superfícies de ataque de crescimento mais acelerado.

A explosão das identidades não humanas

Durante anos, as estratégias de gestão de identidades e acessos (Identity and Access Management, IAM) se concentraram principalmente em usuários humanos. Hoje, essa realidade mudou. As organizações modernas operam milhares de identidades associadas a sistemas automatizados que precisam se autenticar continuamente para acessar dados, serviços e recursos críticos.

Diversos estudos do setor mostram que as identidades não humanas superam amplamente as identidades humanas, em uma proporção média de 50 para 1 em muitos ambientes corporativos. A Okta, uma das empresas líderes em gestão de identidades e acessos, destaca que as organizações modernas administram ecossistemas completos de identidades distribuídas entre usuários, aplicações e cargas de trabalho. Essa expansão aumenta a complexidade operacional e a superfície de ataque.

O que são identidades não humanas?

São credenciais digitais atribuídas a máquinas, aplicações e processos automatizados que costumam operar de forma autônoma e permanente. Em vez de senhas, biometria ou autenticação multifator (Multi-Factor Authentication, MFA) baseada em dispositivos móveis, essas identidades utilizam certificados, chaves de API, tokens e segredos criptográficos para se autenticar e interagir com outros sistemas.

Onde são utilizadas?

Principalmente em cargas de trabalho automatizadas em cloud, como perfis que permitem acesso seguro entre serviços; pipelines de integração contínua e entrega contínua (Continuous Integration/Continuous Delivery, CI/CD) em DevOps que automatizam implantações; comunicação entre microsserviços por meio de OAuth e certificados; integrações de Software as a Service (SaaS) baseadas em APIs; dispositivos IoT conectados; plataformas de análise avançada; e agentes de inteligência artificial que consomem e processam dados automaticamente. Em todos esses casos de uso, as credenciais são necessárias para a operação e se tornam um novo elemento de risco para a cibersegurança.

O risco invisível e em expansão

Cada conta de serviço, token ou certificado representa um possível ponto de entrada para a organização. Como as identidades não humanas costumam operar de forma autônoma, com privilégios elevados e pouca supervisão, o comprometimento de uma identidade privilegiada pode permitir que atacantes acessem dados sensíveis, escalonem privilégios e se movimentem lateralmente no ambiente tecnológico da organização.

Cinco fatores críticos de risco

Excesso de privilégios

- Excesso de privilégios
- Credenciais estáticas de longa duração
- Falta de inventário e visibilidade limitada sobre as identidades existentes
- Gestão deficiente de segredos
- Monitoramento insuficiente de comportamentos automatizados

A inteligência artificial acelera o desafio

A Agentic AI incorpora novas identidades autônomas capazes de consumir APIs, consultar dados e executar ações. Cada agente exige autenticação, autorização e monitoramento contínuo.

Rumo a uma estratégia de Gestão de Identidades Não Humanas (NHIM)

A Gestão de Identidades Não Humanas (Non-Human Identity Management, NHIM) permite descobrir, inventariar, governar, proteger e supervisionar essas identidades ao longo de todo o seu ciclo de vida. Entre as capacidades essenciais estão descoberta automática, menor privilégio, rotação de segredos, monitoramento contínuo, detecção de anomalias e governança unificada.

Recomendações para CISOs

Implementar a descoberta automática de identidades não humanas, aplicar rigorosamente o princípio do menor privilégio, revisar privilégios periodicamente, adotar modelos de Zero Trust, automatizar a gestão de segredos, monitorar comportamentos anômalos e unificar a gestão de identidades humanas e não humanas, atribuindo sempre responsáveis para cada identidade não humana.

Conclusão

As identidades não humanas já não são um tema exclusivamente técnico. São um componente estratégico da resiliência cibernética moderna. As organizações que aplicarem governança adequada a essas identidades estarão melhor posicionadas para reduzir os riscos decorrentes da cloud, da automação, da transformação digital e da inteligência artificial, além de fortalecer sua postura de cibersegurança para proteger seus ativos digitais mais críticos.



César Vega Calderón
Cybersecurity Manager

O risco invisível: Como as identidades não humanas estão redefinindo a cibersegurança

Na era da cloud, da automação e da inteligência artificial, a maioria das identidades que acessam os recursos empresariais não é humana. Protegê-las agora é um requisito estratégico para a resiliência dos negócios.



O CRESCIMENTO QUE NÃO PODEMOS IGNORAR

45:1
Relação média de identidades não humanas por cada identidade humana em ambientes cloud-native.
Fuente: Gartner, 2024

81%
das organizações esperam que as identidades não humanas representem a maioria das identidades nos próximos 3 anos.
Fuente: Okta, 2024

23,8 milhões
de segredos expostos em repositórios públicos do GitHub em 2024.
Fuente: GitGuardian, 2024

98%
dos incidentes relacionados a identidades não humanas poderiam ter sido evitados com visibilidade e controles adequados.
Fuente: CyberArk, 2024

“ A segurança das identidades não humanas é a próxima fronteira crítica. Elas são o elo que mantém nossas aplicações, dados e sistemas conectados. Sem gestão adequada, não podemos proteger o que importa. — Okta Identity Security ”

- Descobrir e visualizar todas as identidades não humanas.
- Aplicar menor privilégio e acesso "just-in-time".
- Gerenciar e rotacionar segredos e certificados de forma automática.
- Monitorar e detectar comportamentos anômalos.
- Governar todo o ciclo de vida com políticas e conformidade.

MENSAGEM-CHAVE
As identidades não humanas são o pilar da transformação digital. Visibilidade, controle e governança são essenciais para transformar um recurso invisível em uma vantagem competitiva.

Segurança sem perímetro: a ascensão da identidade como eixo de proteção

Tendências por Whendy Oré Crisostomo

O modelo tradicional de cibersegurança baseado no perímetro ficou para trás. Em ambientes cloud, distribuídos e altamente dinâmicos, a identidade se consolida como o novo eixo de controle. Essa mudança impacta não apenas os usuários, mas também sistemas, aplicações e processos automatizados. Hoje, proteger significa validar continuamente quem, ou o que, acessa os recursos, inclusive em situações tão cotidianas quanto conectar-se de casa, usar uma aplicação corporativa ou automatizar processos entre sistemas.

A identidade como eixo de controle

A segurança já não depende de estar “dentro” ou “fora” da rede. Cada acesso deve ser avaliado considerando identidade, dispositivo, contexto e nível de risco. Essa abordagem, alinhada ao Zero Trust, elimina a confiança implícita e exige verificação constante. Nesse novo modelo, a identidade se torna o verdadeiro ponto de defesa. Gerenciar identidades de forma precisa e dinâmica é, hoje, a base de qualquer estratégia moderna de segurança.

O crescimento das identidades de máquina

Além dos usuários, o ecossistema digital é dominado por interações entre sistemas. Interfaces de programação de aplicações (Application Programming Interfaces, APIs), microsserviços, containers e automações utilizam identidades de máquina para se autenticar. Essas identidades, como certificados, tokens ou chaves, permitem que aplicações, serviços e processos automatizados se comuniquem entre si sem intervenção humana. Isso ocorre, por exemplo, quando uma aplicação consulta dados de outra ou executa tarefas programadas. No entanto, esse tipo de identidade cresce de forma exponencial e, muitas vezes, carece de visibilidade e controle, tornando-se um ponto crítico de risco dentro das organizações.

Gestão de identidades de máquina: controle e visibilidade

Para mitigar esses riscos, as organizações precisam saber exatamente quais identidades existem, onde são usadas e por quanto tempo permanecem ativas. É nesse ponto que a gestão de identidades de máquina (Machine Identity Management, MIM) ganha relevância.

O objetivo é estabelecer controle integral por meio de:

- Inventário e descoberta contínua
- Automação de credenciais
- Aplicação do princípio do menor privilégio
- Monitoramento constante

Uma estratégia madura de MIM permite reduzir a superfície de ataque e fortalecer a resiliência organizacional.

Por que isso importa agora?

Esse movimento em direção a uma segurança centrada na identidade não é casual. Trata-se de uma resposta à evolução natural dos ambientes digitais, nos quais os usuários já não trabalham a partir de uma única localização e as aplicações já não residem em um único lugar. Hoje, é comum que uma organização utilize ambientes multcloud, aplicações externas e dispositivos pessoais, o que aumenta os pontos de acesso e reduz a eficácia dos controles tradicionais. Nesse contexto, confiar apenas na rede já não é suficiente.

Além disso, os atacantes adaptaram suas estratégias e agora buscam comprometer credenciais em vez de violar infraestruturas diretamente. Isso torna a proteção da identidade, humana e de máquina, um fator decisivo para prevenir acessos indevidos. Não enfrentar esse desafio pode resultar em acessos não autorizados difíceis de detectar, exposição de dados críticos e aumento significativo do risco operacional. Essa mudança reflete uma tendência clara, com a transição de um modelo de segurança baseado em infraestrutura para um modelo centrado na gestão de identidades.

Agentes de identidade: segurança no ponto de execução

Em ambientes distribuídos, os agentes de identidade (ID Agents) levam a segurança diretamente para onde a execução acontece. Esses componentes permitem:

- Validar acessos localmente
- Proteger credenciais sem exposição
- Aplicar políticas em tempo real
- Gerar telemetria de segurança.

A implementação desses agentes reforça o conceito de identidade como perímetro, integrando segurança em cada camada do ambiente digital. Em ambientes modernos, esses agentes permitem que a segurança não dependa de um único ponto central, mas acompanhe cada sistema e aplicação no local em que a execução acontece.

A identidade como um eixo de controle

Mais do que uma evolução tecnológica, essa mudança representa uma transformação na forma como as organizações entendem a confiança. Já não se trata de assumir que tudo o que está dentro do ambiente interno é seguro. O novo modelo exige validar continuamente cada identidade, cada acesso e cada interação. Essa abordagem não apenas melhora a segurança, mas também permite que as organizações operem com mais flexibilidade em ambientes cada vez mais distribuídos.

Conclusão

O perímetro tradicional desapareceu. Em seu lugar, a identidade passa a ocupar o centro da cibersegurança. Compreender essa mudança deixa de ser uma questão exclusivamente técnica e se torna uma necessidade operacional para qualquer organização. Essa evolução não é temporária. Trata-se de uma transformação estrutural na forma como as organizações entendem a segurança.



Whendy Oré Crisostomo
Cybersecurity Analyst



Vulnerabilidades

OneUptime Sandbox Escape

Data: 27 de maio de 2026

CVE: CVE-2026-45102



CVSS: 9,9

CRÍTICA

Descrição

A CVE-2026-45102 é uma vulnerabilidade crítica de execução remota de código (Remote Code Execution, RCE) no OneUptime. A falha está associada ao uso inseguro do módulo "vm" do Node.js como mecanismo de isolamento.

Um atacante autenticado com poucos privilégios pode explorar essa falha para executar código arbitrário no servidor, acessar informações sensíveis e comprometer a aplicação. Devido ao alto impacto, a vulnerabilidade recebeu pontuação CVSS de 9,9.

Solução

- Atualizar para o OneUptime 10.0.98 ou superior, versão em que o problema foi corrigido.

Produtos afetados

- OneUptime: versões anteriores à 10.0.98.

Referências

- nvd.nist.gov
- github.com

Vulnerabilidades

WordPress WP Maps Pro

Data: 29 de maio de 2026
CVE: CVE-2026-8732



CVSS: 9,8
CRÍTICA

Descrição

O WP Maps Pro, plugin do WordPress usado para criar mapas interativos com base no Google Maps, é vulnerável ao escalonamento de privilégios por meio da criação de uma conta de administrador.

Quando o atacante acessa uma página pública, consegue extrair um nonce, ou seja, um valor temporário de uso único exposto no código JavaScript. Em seguida, envia uma requisição AJAX em segundo plano com esse nonce incluído na solicitação.

O sistema não verifica se o usuário tem perfil de administrador. Apenas valida se o nonce está correto, permitindo que o atacante assuma privilégios de administrador.

Solução

- Revisar ou desinstalar plugins instalados recentemente sem autorização.
- Atualizar o plugin para a versão 6.1.1 ou superior.

Produtos afetados

- Versões do WP Maps Pro anteriores à 6.1.0.

Referências

- thehackernews.com
- wordfence.com
- incibe.es

Android corrige vulnerabilidade de escalonamento de privilégios explorada ativamente

Data: 2 de junho de 2026

CVE: CVE-2025-48595

Alta

Descrição

O Google corrigiu a vulnerabilidade CVE-2025-48595, identificada no componente Framework do Android. A falha corresponde a um integer overflow (CWE-190), que pode resultar em elevação local de privilégios e permitir a execução de código com permissões superiores às previstas.

A exploração dessa vulnerabilidade não exige privilégios elevados prévios e pode comprometer a segurança do dispositivo afetado. O Google indicou que havia evidências de exploração limitada e direcionada antes da publicação da atualização de segurança.

A vulnerabilidade afeta múltiplas versões do Android e faz parte do boletim de segurança de junho de 2026.

Produtos afetados

- Android 14
- Android 15
- Android 16
- Android Framework
- Dispositivos Android que não tenham o nível de patch de segurança 2026-06-05 ou superior.

Solução

- Atualizar os dispositivos para o nível de patch de segurança 2026-06-05 ou posterior.

Referências

- android.com
- nist.gov
- cisa.gov

Cisco corrige vulnerabilidade no Unified Communications Manager

Data: 4 de junho de 2026

CVE: CVE-2026-20230

Alta

Descrição

A vulnerabilidade CVE-2026-20230 no Cisco Unified Communications Manager (Unified CM) poderia permitir que um atacante remoto não autenticado realizasse ataques de Server-Side Request Forgery (SSRF), ou falsificação de requisições do lado do servidor, por meio de um dispositivo afetado.

Essa vulnerabilidade está relacionada à validação incorreta de SSRF em determinadas requisições HTTP.

Um atacante poderia explorar a falha enviando uma requisição HTTP manipulada a um dispositivo afetado. Isso permitiria gravar arquivos no sistema operacional de base, que poderiam ser usados posteriormente para obter privilégios de administrador.

Produtos afetados

- Cisco Unified CM até a versão 14;
- Unified CM SME até a versão 15, caso esteja com o serviço WebDialer habilitado.

Solução

- Desabilitar o serviço WebDialer pelos administradores.
- Instalar o patch de segurança mais recente.

Referências

- thehackernews.com
- sec.cloudapps.cisco.com
- incibe.es

Eventos

GITEX AI Europe 2026

30 de junho – 1º de julho

O GITEX AI Europe 2026 será realizado de 30 de junho a 1º de julho, em Berlim, Alemanha, como um espaço dedicado a analisar o papel da inteligência artificial na transformação digital da Europa. O evento reunirá empresas de tecnologia, organizações públicas, centros de pesquisa e líderes do setor para compartilhar experiências sobre a adoção da IA, os desafios regulatórios e as oportunidades de inovação que essa tecnologia está gerando em diferentes setores. A programação incluirá conferências, demonstrações e debates sobre o desenvolvimento de soluções baseadas em IA, assim como seu impacto em áreas como cibersegurança, produtividade empresarial e competitividade digital.

[Enlace](#)

RAISE Summit 2026

7 a 9 de julho

O RAISE Summit 2026 será realizado de 7 a 9 de julho, em Paris, França, reunindo milhares de líderes empresariais, investidores, empreendedores e responsáveis por tecnologia para debater o futuro da inteligência artificial. O evento contará com a participação de mais de 2.000 empresas e centenas de palestrantes de referência do ecossistema tecnológico global, incluindo representantes de organizações como Google, OpenAI, Anthropic, Meta e AWS. Ao longo da programação, serão analisados os desafios associados à adoção empresarial da IA, desde infraestrutura e governança até geração de valor para o negócio, consolidando o evento como um espaço relevante para conhecer as estratégias que estão definindo a próxima geração da inovação digital.

[Enlace](#)

Gartner Security & Risk Management Summit

22 a 24 de julho

O Gartner Security & Risk Management Summit será realizado de 22 a 24 de julho, em Tóquio, e reunirá CISOs, responsáveis por risco e líderes de cibersegurança para analisar os desafios que estão redefinindo a função de segurança nas organizações. A programação será estruturada em seis trilhas especializadas: liderança e inovação em cibersegurança, gestão de riscos e resiliência, segurança de infraestruturas e cloud, segurança de aplicações e dados, operações e resposta a incidentes, além do programa exclusivo CISO Circle para executivos. Nesses espaços, os participantes poderão conhecer as principais tendências em inteligência artificial, resiliência cibernética, proteção de dados, segurança cloud e alinhamento das estratégias de segurança aos objetivos de negócio.

[Enlace](#)

Recursos

➤ ENISA NIS360

Publicado pela Agência da União Europeia para a Cibersegurança (ENISA), o relatório NIS360 2026 oferece uma avaliação abrangente do nível de maturidade e criticidade em cibersegurança dos setores considerados essenciais pela Diretiva NIS2. O estudo analisa o nível de prontidão de setores-chave como energia, transporte, finanças, saúde, infraestrutura digital e administração pública, identificando pontos fortes, lacunas e áreas prioritárias de melhoria. Além disso, oferece uma visão comparativa da evolução da resiliência cibernética na Europa e destaca os setores em que a importância estratégica supera o nível atual de maturidade em segurança, tornando-se uma referência importante para responsáveis por cibersegurança, reguladores e operadores de infraestruturas críticas.

[Enlace](#)

➤ NIST IR 8374r1: Ransomware Risk Management – A Cybersecurity Framework 2.0 Community Profile

Publicado pelo National Institute of Standards and Technology (NIST), este documento oferece uma orientação prática para ajudar as organizações a gerenciar o risco associado ao ransomware usando o NIST Cybersecurity Framework (CSF) 2.0. O perfil identifica os objetivos de segurança mais relevantes para governar, prevenir, detectar, responder e se recuperar de incidentes de ransomware, fornecendo uma referência estruturada para avaliar o nível de prontidão de uma organização diante desse tipo de ameaça. Além disso, o documento facilita a identificação de lacunas de segurança e a definição de estratégias de resiliência que permitam reduzir o impacto operacional e financeiro de um ataque.

[Enlace](#)



**CLIQUE AQUI E
SIGA A RADAR**

**Powered by the
cybersecurity
NTT DATA team**

br.nttdata.com