

The Licensing Journal, Key Issues in Agentic AI Implementation and Integration Deals, (Aug. 1, 2026)

The Licensing Journal

Key Issues in Agentic AI Implementation and Integration Deals

Rohith George, Brad Peterson, Arsen Kourinian, Joe Pennell, Veronica Glick, Nitya Sunil and Oliver Yaros

Rohith George is based in Mayer Brown's San Francisco office and Co-Leader of its Technology & IP Transactions practice. He advises companies seeking to buy, sell, or build technology products and services, including advising on contracts involving AI, digital infrastructure, complex outsourcing, cloud solutions, embedded financial services, and other emerging technologies.

Brad Peterson is a Mayer Brown partner in the Chicago office and a member of the Technology & IP Transactions practice. Brad's practice is currently focused on helping clients negotiate cloud, systems integration and outsourcing deals that reduce cost, improve operations and accelerate adoption of AI technologies.

Arsen Kourinian is a partner in Los Angeles in Mayer Brown's Cybersecurity & Data Privacy and Artificial Intelligence practice groups who counsels clients regarding compliance with data privacy, cybersecurity and artificial intelligence laws, processes for building an artificial intelligence governance program, and addressing these regulatory requirements in transactions and other corporate deals.

Joe Pennell is a Chicago-based partner and Co-Leader of Mayer Brown's Technology and IP Transactions group. He advises clients contracting for AI, digital infrastructure, transformation initiatives, data center leasing and services, cloud and SaaS arrangements, software licensing and implementation, technology collaborations, and complex outsourcing transactions.

Veronica Glick is a Cybersecurity & Data Privacy partner who advises clients on cybersecurity, national security, emerging technology, artificial intelligence, cyber incident response, and security-related M&A diligence.

Nitya Sunil is an associate in Mayer Brown's Chicago office, counseling clients on complex commercial transactions involving technology and business process outsourcing, technology licensing, and intellectual property rights.

Oliver Yaros is the leader of Mayer Brown's London office Intellectual Property group, who advises clients on technology and outsourcing transactions with a particular focus on artificial intelligence, fintech, and digital transformation projects, as well as data protection, cybersecurity incidents, and intellectual property matters.

Companies are engaging systems integrators to design, build, and operate agentic artificial intelligence (AI) solutions. The integration agreements for those engagements raise contract issues distinct from traditional system integration, including AI liability allocation, IP rights in AI-generated work product, vendor lock-in, and governance of autonomous agents acting at speed and scale. This article covers critical issues including deal structure, performance SLAs, dependency management, technology stack selection, AI governance, data privacy and security, IP ownership, and exit planning. [\[1\]](#)

PATHS TO ADOPTING AGENTIC AI

Which AI Adoption Model Fits Your Enterprise: DIY, Outsourcing, or Integration?

Companies today have multiple paths to adopting agentic artificial intelligence (AI). On the do-it-yourself path, the company procures AI tools directly and uses its own technology, legal, compliance, and business teams to adapt, configure, orchestrate, and govern them. That path provides control but is limited by the company's capabilities. On the outsourced-function path, a service provider performs a business process for the company and uses agentic AI within the provider's own delivery model to improve efficiency or reduce cost. That involves some loss of control, but the company can rely on the provider's capabilities and commitments to deliver business outcomes.

This article describes a third path. On that path, the company engages a consultant, systems implementer, managed service provider, or other integrator to design, build, integrate, and support AI agents within the company's environment. The company maintains control through its governance of the overall project but leverages the skills of the integrator to create a working, integrated solution from the company's existing code stack, the integrator's proprietary code, and a variety of third-party AI tools.

Leading agentic AI companies are building alliances with major consulting and technology firms to move enterprise customers along that path from pilots to production. Those alliances reflect the practical reality that deploying agentic AI at scale often requires third-party expertise in business process redesign, systems integration, data engineering, security review, change management, operating-model decisions, and ongoing monitoring.

What Should the Scope of Work Cover in an AI Integration Contract?

The integrator's scope of work should start with the affected business processes, not merely with the agents to be built. In many projects, the most important deliverable is a redesigned business process that specifies which tasks will be automated, which decisions the agent may make, where human review remains required, what controls apply, and how the company will adopt the new way of working. That scope may range from a narrow implementation of a point solution to a broad AI-enabled transformation of the company's operating model. Some possible components include:

- **Advise/Design:** The integrator evaluates the company's operations and recommends where to deploy agentic AI, how to redesign the affected business processes around agents, which AI tools to use, and/or what the target operating model looks like. The integrator may also develop the business case and propose a change management approach for the project.
- **Build/Integrate:** The integrator helps the company procure AI tools, integrates them into the company's technology environment, builds workflows, configures agents, connects systems and data sources, and implements agreed controls.
- **Operate/Support:** The integrator may provide primary and secondary support. With primary support, the integrator operates, monitors, and manages the agentic AI solutions to ensure they function as designed and adapts them as business needs evolve. With secondary support, the integrator provides training and support for escalated issues as required to allow the company or its designee to provide primary support for the integrated solutions.

Each agentic solution will be designed, built, and then supported. However, the integration likely will involve the gradual growth of a network of AI agents, workflows, and integrations operating across the company's systems, with some AI agents being supported as others are designed and built.

This article addresses only the contract with the integrator as an integrator. Different provisions would apply if the integrator also provides or hosts the AI tools or provides a managed service based on the integrated agentic AI

solutions. For this article, we assume that the company contracts directly with third parties for the agentic AI tools, which are hosted by the company or its agentic AI provider, and that the integrator is authorized to work with the agentic AI tools, as hosted, to provide services to the company under the third-party agreements. That structure is attractive to many companies, in large part because internal privacy and cybersecurity requirements favor enterprise versions of AI platforms with inputs and outputs remaining within the company environments.

Critical Contract Issues

These critical contract issues face companies considering contracting for integrated agentic AI solutions:

- Deal structure;
- Performance commitments and dependencies;
- Technology and stack selection;
- Project governance;
- Liability;
- Security, privacy, and data use;
- IP rights; and
- Exit and transition.

These issues are in addition to the other issues seen in commercial contracts generally and those specifically for consulting and other systems integration.

How Should You Structure an Agentic AI Integration Contract?

As is often the case in integration agreements, the fundamental challenge is that the contract will be signed well before key cost drivers and operational and technical constraints are known, and critical decisions are made. This uncertainty is even more pronounced with agentic AI programs because the tools are evolving quickly, the target workflows may change during discovery, and the ultimate solution may depend on multiple third-party providers, an internal IT function, and business units outside the integrator's scope that nonetheless depend on shared IT infrastructure and AI systems.

As a result, before completing design, the integrator may have difficulty estimating how much it will cost or whether it is even possible to achieve the company's specified automation targets, cost savings, or other desired business outcomes. At the same time, the company may be reluctant to fund a broad design and integration effort without some confidence in the eventual cost and value proposition.

These challenges are addressed with deal structures. Most deal structures fall into three categories—"assist," "deliver," and "shared." A single program may use different structures at different phases. For example, the parties might use an assist model for the design phase (where scope is uncertain) and transition to a deliver or shared model for the build and support phases (where requirements are better defined).

Assist

In an assist deal structure, the integrator works at the company's direction without committing to outcomes, typically paid on a time-and-materials basis—although assist-phase work with a defined scope, such as innovation workshops analyzing an identified set of the company's business processes, may be priced on a fixed-fee basis. This approach allows the parties to start quickly and adapt the work as the facts develop. It is often appropriate for early assessments, use-case identification, and design work where the parties do not yet know what should be built.

The tradeoff is that the company bears the risk of budget overruns and schedule delays. The integrator has a limited downside if the project expands, so the company must, through governance, manage the cost of the project. A well-crafted contract can allow the company to mitigate that risk through budget controls, staffing commitments, reporting obligations, and change-control mechanics that allow the company to make informed decisions before costs run ahead of value, but the company will need the knowledge and skill to exercise that control.

Deliver

In a deliver deal structure, the company agrees to pay the integrator a fixed fee or milestone-based price. This can give the company price certainty for a clearly defined scope, and it gives the integrator a strong incentive to complete the work efficiently.

The benefits depend on the company precisely defining its desired end results—a challenge, given that requirements may evolve during the project. Agentic AI solutions may operate across disparate platforms, environments, business units, functions, and processes (with different IT owners, buyers, and users), all of which will be changing. The company may also still be developing or evolving its comprehensive enterprise-wide AI or agentic strategy.

If the desired end results aren't defined clearly and accurately, the parties may spend the project negotiating change orders. or in disputes. That risk is heightened by the integrator's financial incentive to understate scope and price to win the deal and to upsell after the deal is won and the company's financial incentive to increase benefits without increasing cost. The integrator will also charge a risk premium, especially if it is being asked to commit to automation, savings, or production performance before it has completed discovery and without firm commitments by the company as to the project's scope and the company's obligations to change its operations as needed for the project.

Shared

In a shared-risk or shared-reward deal structure, the focus is on sharing overall risk by aligning incentives on cost, speed, or measured business outcomes. The contract might establish a target budget, allocate control over major cost drivers, permit charges up to the target budget at agreed rates, share underruns, reduce rates after the target is exceeded, provide bonuses for early delivery, or tie compensation to agreed productivity metrics. For example, if the company determined that every second of reduction in average cycle time for a process would be worth \$100 per month, the contract could provide the integrator with a bonus of \$50 per month for some period for each second saved.

A shared model can be attractive where the company wants the integrator to have a stake in the outcome. But, like the deliver model, it requires an up-front investment in carefully defining desired outcomes and addressing changes in scope. In particular, the contract should define the baseline, measurement period, data sources, re-baselining triggers, company dependencies, chargeable and non-chargeable changes, and audit rights. A benefit of the shared model is in allowing more flexibility in those definitions, recognizing that the parties can jointly affect them, and thus providing additional ways to address uncertainty.

What Performance Commitments Belong in AI integration Contracts?

Performance commitments in agentic AI integration agreements should be separated into at least four categories: implementation commitments, operational commitments, agent-performance commitments, and business-outcome commitments. The further the contract moves from implementation deliverables toward business outcomes, the more important it becomes to define assumptions, dependencies, and relief mechanisms. In framing these commitments in the contract, drafters should keep in mind that the company's objective is rarely an agent in isolation. Instead, it is a redesigned business process that completes work faster, better, or with fewer people. Contracts that define success only at the agent level may deliver working agents without delivering value.

- Implementation commitments address deliverables, milestones, acceptance criteria, documentation, training, and deployment dates.
- Operational commitments address support, monitoring, incident response, uptime, latency, escalation, drift management, and ongoing maintenance.
- Agent-performance commitments address the agents' success, autonomy, quality, and reliability through metrics like completion rate, human-handoff rate, rework rate, and tool-use reliability rate.
- Business-outcome commitments address cycle-time reduction, cost reduction, error reduction, customer experience, productivity, capacity creation, or workforce impacts.

In addition to defining required performance, the contract should define how that performance will be tested. Because agentic AI outputs are non-deterministic, a single successful test is not enough to demonstrate that a solution will perform as required in production. Instead, the contract might require acceptance testing based on statistical performance against an agreed evaluation set, with defined thresholds for accuracy and latency. Testing of business outcomes should not just assess how the solution functions in a development environment, but how it performs as part of the ongoing business operations over time.

The integrator will reasonably argue that its ability to deliver results depends heavily on the company's actions. For example, the integrator depends on the company to approve recommendations, provide access to systems and data, invest in structuring related business data and making necessary changes to its systems and business processes, make SMEs available, participate in testing, make policy decisions, drive the needed organizational change and adoption, not block or delay integration, and, where cost takeout is part of the business case, make workforce or redeployment decisions. The contract should identify these dependencies with enough specificity that the parties can determine whether a missed target was caused by the integrator, the company, an upstream provider, or a shared risk.

Relief should not be automatic or unlimited. The integrator should be required to notify the company promptly of a failed dependency, explain the effect of the failed dependency, mitigate where commercially reasonable, continue to perform the unaffected work, and complete a root-cause analysis to verify that the failed dependency caused the performance failure.

Who Controls the AI Technology Stack in an Integration Deal?

An integrator-led agentic AI solution may be built from a mix of first- and third-party technology, including foundation models, agent platforms, orchestration tools, databases, enterprise applications, monitoring tools, and custom code. The integrator and the company both have good reasons to want to select the "stack" of products being integrated. The integrator will prefer tools that it knows well, can integrate efficiently, can support profitably, provide experience that can be used to win future business, are consistent with agreed performance requirements, and align with its

alliance relationships. Companies will focus on total cost of ownership, higher reliability, greater security, regulatory posture, lock-in risk, and alignment with their enterprise architecture. The company should also consider whether the integrator has incentives to prefer a tool because of the integrator's alliance or other financial relationship with the provider of that tool.

If the integrator selects a tool, the company will expect the integrator to take more responsibility for the performance of that tool within the integrated solution. The integrator will resist guaranteeing upstream tools it does not control and prefer to contractually characterize the choices as recommendations that the company has decided to accept. That reduces the value of a "turnkey" proposal. Even where upstream outages or other failures are carved out, the integrator may still be responsible for designing fallback processes, substitution plans, monitoring, and incident coordination.

A possible approach is to give the integrator discretion to select tools that meet pre-approved criteria, such as (a) platforms listed in the contract; (b) tools already approved for the company's internal use; and (c) tools satisfying the company's security, regulatory, privacy, data use, and performance requirements. Use of tools outside those criteria would require company approval, with defined turnaround times, standards for acceptance, escalation procedures, and effects on liability if the integrator proceeds without approval. The contract should also address migration rights in the event of material changes in third-party terms, prices, functionality, data use, security posture, or model performance.

What Governance Frameworks Do Agentic AI Integration Projects Require?

The functional specifications will almost certainly be incomplete when the contract is signed. To address that problem, the contract should include a robust governance framework for the integration project. Equally important, the underlying tools and platforms will continue to evolve after signing, so the governance framework must be designed to flex with them rather than lock in a point-in-time architecture. For example:

- **Policy and risk governance** is generally controlled by the company and implemented with the support and advice of the integrator. Policy decisions might include defining which categories of data the AI agents may process, which systems they may access, what actions they may take, what risk tolerances apply, and what legal or regulatory requirements govern the use cases. The parties would each commit to defined turnaround times for urgent policy changes, such as those involving regulatory, legal, security, or other business risk.
- **Technical change management governance** is generally a joint process for material changes to model versions, prompts, workflows, orchestration logic, guardrails, escalation rules, integrations, agent permissions, and decision logic. If experience shows that an agent is producing unacceptable outputs or decisions, the governance process should identify who can suspend the AI agent, change its authority, require additional human review, or roll back the release.
- **Audit, observability, and reporting** rights generally run to the company and are implemented in the agentic AI solution by the integrator. For agents acting autonomously or semi-autonomously at speed and scale, observability is the oversight mechanism. By observability, we mean a set of activities designed to confirm that agents are operating within their defined scope of authority and acceptable range of performance. Observability may include monitoring, validating, sampling, and auditing the agents' decisions, actions, and outputs on an ongoing basis, with the contract expressly allocating responsibility for each among the company, the integrator, and any operator of the business process. The company should have access to, for example, agent logs, decision traces, tool-use records, prompts and system instructions, model and version identifiers, evaluation results, and drift indicators, on terms sufficient to support internal oversight, regulator inquiries, and incident investigation. The contract should define retention periods, format, and delivery mechanics (e.g., read-only

dashboards, scheduled exports, or API access), and should preserve these rights through termination and for a reasonable tail period.

- **Day-to-day operational governance** is generally managed by the company, but may be managed by the integrator during the support period with an obligation to stay within agreed parameters. This would include adapting to upgrades and changes by third parties, break/fix and other incident management, and ordinary course testing and maintenance to keep the agents working within defined parameters.

Who Is Liable When AI Agents Cause Harm?

When harm arises from an agentic AI system, the relevant failure may arise from how the integrator selected, configured, integrated, governed, or operated the AI tools. Alternatively, it may arise from parts of the solution required or provided by the company (including its agentic AI tool providers) or from how the company deployed and used the integrated system in production.

A useful principle is to align responsibility with control. The party with the right and practical ability to make a decision should generally bear more risk for the consequences of that decision. If the company dictates use of a particular AI provider against the integrator's recommendation, the integrator will seek relief for failures caused by that decision. If the integrator selects the stack, configures the agents, designs the workflow, and implements its selections in the company's environment, the company will expect the integrator to stand behind the integrated solution, at least where the failure could have been prevented through reasonable design, controls, monitoring, fallback processes, or incident response.

In practice, that principle will often be difficult to apply because many causes of failure sit outside any one party's complete control. Agent behavior may be affected by company data quality, undocumented process rules, user adoption, third-party model changes, enterprise-system outages, security constraints, or post-go-live operating decisions. Integrators, therefore, will often resist broad guarantees of agent behavior, business outcomes, or downstream legal compliance unless the contract clearly defines assumptions, dependencies, authority boundaries, monitoring obligations, relief events, and the scope of post-go-live responsibility.

Caps and exclusions require special attention. Integration agreements often waive consequential and indirect damages and cap remaining liability by reference to fees, with exceptions and "super caps" for certain categories. That structure may be difficult to calibrate when the integrator's fees may be modest relative to the potential business, regulatory, or security harm caused by autonomous AI agents acting at speed and scale. Neither party can model the risk well. The technology is new and difficult to test because of the variability of generative AI systems underlying and guiding AI agents. The result may be that the liability negotiations drive the deal structure, at least until market standards develop in this area.

Indemnification should be addressed separately from liability caps. Traditional IP indemnities may not cover claims arising from actions of AI agents, including claims that their output infringes third-party intellectual property, that their decisions violate applicable laws or that their processing of data breached privacy or confidentiality obligations. The parties should address third-party claims arising from the integrator's selection, configuration, integration, operation, or governance of the AI tools; claims arising from company-directed uses; and, because the indemnities from AI providers are typically narrow and heavily conditioned, the gaps between upstream AI-provider indemnities and the company's likely exposure and whether and how those gaps are absorbed, insured, or otherwise mitigated.

Who Is Responsible for Compliance with Changing AI laws?

Agentic AI integration contracts should allocate responsibility for identifying, interpreting, and operationalizing applicable AI laws. The allocation may depend, for example, on the project model, the use case, the deployment geography, and the parties' legal roles, including whether a party is acting as a developer/provider, deployer, operator, or support provider.

The contract should require a compliance workstream that classifies use cases, identifies applicable legal requirements, assigns responsibility for impact assessments and documentation, and defines how regulatory changes will be handled through governance and change control. The parties should also address whether new legal requirements are included in the existing fees and commitments, treated as chargeable changes, or handled through a shared regulatory-change mechanism.

How Do You Address Data Privacy and Security in AI Integration Contracts?

Agentic AI systems often require data and system access that is broader, more dynamic, and more operationally sensitive than a typical software implementation. AI agents may access enterprise applications, retrieve documents, call APIs, write to systems of record, generate logs, and transmit prompts, outputs, or telemetry outside the company's systems, whether to cloud-based AI models, other third-party providers, or the integrator's own environment.

Generally, the company will, as data controller, retain the rights to define what categories of data may be accessed by the AI agents, where it may be processed, by what companies, for what purposes, and under what security, as well as retention and deletion controls. The integrator would agree to comply with the Company's direction in those areas.

With respect to technology provided by the integrator or its suppliers, the contract should specify whether company data, prompts, outputs, logs, embeddings, telemetry, evaluation sets, and derived artifacts may be used to train or improve the integrator's (or its suppliers') models or other services. If the expectation is "no training," the contract should say so expressly and should require the integrator to select provider tiers, configurations, and sub-processors consistent with that commitment.

Agentic AI systems present an expanded attack surface as agents can autonomously traverse systems and take consequential actions across the enterprise with limited human intervention. Security controls must therefore be agent-specific, rather than generic. Appropriate controls will vary by deployment context. For example, human review may be essential for decisions affecting HR data flows, yet counterproductive when applied to Endpoint Detection and Response tooling where speed is a core performance requirement.

The contract should address evolving security standards across several domains: identity and access management (e.g., least-privilege access, credential management, segregation of duties); operational controls (approval gates for privileged actions, tool-use logging, human review thresholds for high-risk actions); data governance (prompt and output retention, data localization, cross-border access); security assessments (testing resilience to AI attack vectors such as prompt injections and data poisoning); and security protocols (monitoring for anomalous behavior, incident response, vulnerability management, and subcontractor controls). For regulated use cases, the contract should also allocate responsibility for DPIAs, AI impact assessments, model risk assessments, audit documentation, and regulatory-response support.

Who Owns AI-Generated Work Product in an Integration Deal?

Each party will usually want to retain and reuse its background IP, methodologies, templates, connectors, prompts, and know-how, and may even have obligations to do so under its partnership or alliance agreements with AI providers or other third-party agreements. The company should consider whether any project work product provides a strategic

advantage, embeds company confidential information, or could create a vulnerability if reused for another of the integrator's customers. If so, the company may have a compelling argument for restrictions. If not, the integrator may prevail on the logic that its business depends on bringing IP from prior engagements and leaving with improved IP. In either case, the contract should define clear rights to background IP, company-specific deliverables, custom code, configurations, prompts, evaluation sets, fine-tuned weights, and other artifacts created for the company.

A newer set of questions concerns intellectual property rights in artifacts, code, decisions, content, recommendations, and other outputs generated by AI agents. Standard work-product assignment clauses drafted for consulting and systems integration engagements may not map cleanly onto outputs from AI agents because those clauses often assume human authorship. If the integrator is intending to use AI agents in the performance of its own services for the company, this creates an ambiguity about whether outputs from those agents are IP owned by the integrator or the company. Thus, the contract should expressly address, as between the company and the integrator, who may use, disclose, modify, retain, and commercialize those outputs, including outputs derived from company data or confidential information.

Creating and protecting rights in these outputs requires additional provisions and effort. Current U.S. Copyright Office guidance and recent case law indicate that material generated by AI without sufficient human authorship is not protectable by copyright, and courts have declined to recognize AI solutions as patent inventors. A traditional assignment of "all right, title, and interest" in output may convey less than the company expects because copyright and patent rights may never exist for material generated solely by AI. If the parties wish to seek trade secret protection for agentic outputs, the contracting might require confidentiality controls, access restrictions, marking practices, and retention rules required to support the company's protection strategy. Similarly, if the company wishes to seek copyright protection for agentic outputs, the contracting might require human-in-the-loop steps and contemporaneous records of human contribution.

The integrator's agentic output also might infringe third-party rights or introduce non-compliant open source into the company's proprietary stack. The parties should address who bears that risk, whether the integrator provides an indemnity (and on what conditions), whether upstream AI provider indemnities are passed through, and what operational controls the integrator must implement, such as output scanning, code scanning, license filtering, restricted use cases, or human review.

What Exit Rights Should Enterprises Negotiate in AI Integration Contracts?

Exit and transition are as important in agentic AI integration agreements as in traditional IT and business process outsourcing. The mechanics may be different because operational knowledge may reside in prompts, configurations, orchestration logic, evaluation sets, monitoring rules, runbooks, and integration documentation. Without clear terms, the company may discover at exit that it does not have what it needs to operate the agents itself or with another provider.

The contract should define what the company receives and what the integrator must do upon any expiration or termination, including: the scope and duration of transition assistance; delivery of agent configurations, prompts and system instructions, orchestration logic, custom code, evaluation sets, documentation, runbooks, support procedures, and training materials; and knowledge transfer sufficient to allow the company or a successor to operate the solution.

The company should also have rights to continue receiving services during transition, export data and logs in usable formats, receive reasonable assistance from the integrator, and retain necessary license rights for a transition period. These rights should be fully negotiated at signing, when the company has leverage, rather than after the integrator has become embedded in the operating model.

TAKEAWAYS

- Agentic AI integration agreements are more complex than traditional enterprise system integration contracts because they blend design, build, and support around a technology stack that changes in real time.
- Effective AI integration contracts must balance flexibility for rapid changes in AI capabilities with enforceable structure, accountability, and risk protection for the enterprise.
- Key provisions must address AI liability gaps, performance SLAs, technology stack selection, portability of prompts and configurations, IP rights in AI-generated work product, and transition assistance at exit.
- Companies that negotiate clear deal structures and governance frameworks will capture the benefits of agentic AI without outsized liability exposure or vendor lock-in.

Whether contracting directly with AI providers or through a systems integrator, the core objectives—defining scope, establishing AI governance guardrails, creating accountability for AI agent behavior, and allocating risk for autonomous AI harm—require new contracting approaches tailored to agentic AI's unique characteristics.

Footnotes

- ¹ See other articles in this series including “Contracting for Agentic AI Solutions: Shifting the Model from SaaS to Services,” at <https://www.mayerbrown.com/en/insights/publications/2026/02/contracting-for-agentic-ai-solutions-shifting-the-model-from-saas-to-services>

[https://prod.resource.cch.com/resource/scion/document/default/\(%40%40PAA01%20WK-JSTORY20260801-4\)paa0110c86fb1c01744fb831a6188d7c83eb8?cfu=Legal&cpid=WKUS-Legal-Cheetah&uAppCtx=cheetah](https://prod.resource.cch.com/resource/scion/document/default/(%40%40PAA01%20WK-JSTORY20260801-4)paa0110c86fb1c01744fb831a6188d7c83eb8?cfu=Legal&cpid=WKUS-Legal-Cheetah&uAppCtx=cheetah)