

THE FUTURE BELONGS TO AGENT ORCHESTRATORS

Building and Managing the AI Workforce That Will Redefine
Business Leadership and Personal Productivity

THERELEE D. WASHINGTON II

Technology Leader Entrepreneur Educator AI Strategist

Executive Summary

Artificial intelligence is moving from a tool people consult to a workforce they direct. The decisive advantage will not belong merely to those who know how to prompt a model. It will belong to people and organizations that can design, deploy, govern, and continuously improve collections of specialized AI agents while preserving human judgment and accountability.

An AI agent is more than a chatbot. It is a software-based worker configured around an objective, equipped with relevant knowledge and tools, and permitted to take defined actions within guardrails. A single agent may research, analyze, draft, monitor, classify, coordinate, or execute. An agent portfolio connects these capabilities to business and personal objectives, with humans setting direction, approving consequential decisions, and measuring results.

My own work increasingly operates through this model. Across an educational institution, a technology consulting company, professional-services responsibilities, public thought leadership, and personal planning, I use specialized AI workspaces and agents to preserve context, investigate problems, prepare decisions, develop content, monitor work, and convert strategy into execution. The value is not that AI replaces leadership. The value is that it gives leadership greater reach—provided the system is intentionally designed and responsibly managed.

This paper introduces the AGENT Portfolio Framework: Aim, Guardrails, Expertise, Network, and Tracking. It provides leaders with a practical method for deciding what an agent should accomplish, what authority it should possess, what knowledge and tools it needs, how it coordinates with people and other systems, and how its performance and risk will be measured.

The central conclusion is straightforward: the future of work will be shaped by agent orchestrators. These will be people who can translate objectives into agent roles, integrate agents into real workflows, develop quality standards, supervise exceptions, and create productive partnerships between human talent and digital intelligence. This is not only a technical capability. It is a new discipline of leadership.

The Argument in Brief

Shift	Old operating assumption	Agent powered operating assumption
Use of AI	An employee asks a general assistant for help	A portfolio of specialized agents advances defined objectives
Management	Leaders assign work only to people	Leaders coordinate people automation software and agents

Productivity	Faster completion of isolated tasks	Redesign of complete workflows and decision cycles
Quality	Review the final answer informally	Evaluate outputs against standards evidence and outcomes
Risk	Control access to a single tool	Govern identity permissions data actions and agent interactions
Advantage	Access to a capable model	Superior context orchestration governance and learning

The End of AI as a Single Assistant

The first phase of generative AI trained millions of people to think in a simple pattern: open a chat window, ask a question, receive an answer, and repeat. That pattern remains useful, but it does not constitute an operating model. It depends on the user to restate context, recognize every next step, move information between systems, and judge quality without a consistent standard.

The agent era changes the unit of work. Instead of asking what one model can say, we ask what a configured system can accomplish. Current agent platforms define agents as applications that can plan, call tools, collaborate across specialized roles, and maintain enough state to complete multi-step work. The difference is agency: an agent can work toward an objective across a sequence of decisions rather than merely produce one response.

This does not mean every workflow should become autonomous. A deterministic process remains the right choice when the rules are stable and every step is known in advance. An agent becomes valuable when work contains ambiguity, unstructured information, exceptions, judgment, or changing paths that are difficult to encode as traditional rules. The most effective organizations will know when to use a person, a conventional automation, an agent, or a combination of all three.

The practical shift is from isolated assistance to coordinated capability. A research agent may gather evidence. A strategy agent may test alternatives. An operations agent may translate a decision into tasks. A communications agent may prepare audience-specific messages. A quality agent may challenge assumptions and check the work against policy. A human leader may review the evidence, resolve tradeoffs, and authorize action.

The future is therefore not one super-agent operating a business without oversight. It is a managed network of focused capabilities, each assigned an appropriate role and level of authority.

From AI Users to Agent Orchestrators

The labor-market distinction emerging around AI is often described as the difference between people who use AI and people who do not. That distinction is already becoming too shallow. Soon, access to a capable assistant will be ordinary. The larger performance gap will be between people who use AI episodically and people who organize AI systematically.

Microsoft's 2025 Work Trend Index reported that 28 percent of managers were considering hiring AI workforce managers and 32 percent planned to hire AI agent specialists within 12 to 18 months. Leaders expected their teams to redesign processes with AI, build multi-agent systems, train agents, and manage them. Its 2026 findings further connected effective AI adoption to managers who model use, establish quality standards, encourage experimentation, and reward work redesign. These findings describe a management transition, not simply a software rollout.

An agent orchestrator performs five leadership functions:

- Translates strategic objectives into clearly bounded agent responsibilities.
- Selects the appropriate combination of human expertise, automation, models, data, and tools.
- Designs handoffs among agents and between agents and people.
- Establishes quality, safety, privacy, and approval standards.
- Measures results and improves the system based on evidence rather than novelty.

These capabilities can be practiced by a corporate executive, small-business owner, educator, consultant, nonprofit leader, ministry leader, student, or household decision-maker.

Orchestration is not reserved for software engineers. Technical professionals will build many of the underlying systems, but operational leaders must define the work, risk, value, and human responsibilities.

The strongest agent orchestrator is not necessarily the person who writes the most complex prompt. It is the person who understands the objective, sees the whole workflow, identifies what can go wrong, provides the right context, and creates a reliable system for producing and reviewing results.

What an AI Agent Actually Is

The word agent is increasingly applied to everything from a saved prompt to fully autonomous software. A useful operational definition must be more precise.

An AI agent is a goal-directed system that uses an AI model to interpret context, make decisions, select or use tools, and complete one or more steps on behalf of a person or organization within defined boundaries.

Seven components distinguish a functioning agent from a casual chat:

Component	Purpose
Objective	Defines the outcome the agent is responsible for advancing.
Role and instructions	Establishes responsibilities methods limits and expected behavior.
Context	Provides organizational history policies definitions preferences and current state.
Knowledge	Connects the agent to trusted documents records databases or authoritative sources.
Tools	Allows the agent to search calculate create files update systems send requests or perform other actions.
Memory and state	Preserves what has happened and what remains to be done across steps.
Guardrails and evaluation	Limits authority and tests whether results meet accuracy safety and business standards.

Agents can operate at different levels of autonomy. At the lowest level, an agent prepares work for human review. At a moderate level, it may execute routine actions and escalate exceptions. At a high level, it may plan and perform multi-step work with periodic human checkpoints. Authority should increase only after the agent demonstrates reliable performance under controlled conditions.

The Agent Portfolio

A business rarely succeeds by hiring one person to perform every function. The same principle applies to agents. A single general agent asked to understand every policy, customer, workflow, technical domain, risk, and communication style will become difficult to control and evaluate.

A portfolio approach creates specialized responsibilities. Each agent has an explicit purpose, a defined source of truth, a limited set of tools, and known escalation points. The portfolio is then governed as a system.

A practical portfolio may include:

Agent category	Representative responsibilities	Typical human owner
Executive intelligence	Strategic briefings scenario analysis decision preparation	Executive or founder
Operations	Work queues status monitoring handoffs documentation	Operations leader
Customer or student engagement	Triage guidance follow up and escalation	Service or student success leader
Revenue and growth	Lead research opportunity qualification campaign analysis	Sales or marketing leader
Knowledge and compliance	Policy retrieval evidence checks deadlines and exception flags	Compliance or legal owner
Content and communications	Drafting adaptation publishing preparation brand consistency	Communications leader
Quality assurance	Fact checks rubric scoring policy review adversarial challenge	Independent reviewer

The portfolio must be connected to priorities. An impressive collection of agents that does not reduce cycle time, improve decisions, increase service capacity, protect quality, or create measurable value is simply another layer of technology to manage.

My Journey Toward an Agent Powered Operating Model

I did not begin with the goal of building an AI workforce. I began with a familiar leadership problem: the number and complexity of objectives exceeded the hours available to address them well. I was simultaneously leading businesses, supporting an educational institution, managing demanding professional responsibilities, developing new initiatives, producing public thought leadership, and navigating important personal objectives.

The earliest value came from using AI to accelerate research and writing. The deeper transformation occurred when I stopped treating every conversation as a new interaction and began creating persistent, purpose-specific operating environments. Each environment held the relevant history, terminology, standards, documents, decisions, and unfinished work for a domain. Over time, these environments began to function less like notebooks and more like specialized members of a leadership team.

I created distinct capabilities for institutional strategy, recruitment and enrollment, educational operations, program development, compliance, financial analysis, technology solution design, professional-services scoping, publishing, website development, career strategy, and personal planning. The goal was not to eliminate people. It was to ensure that important context was preserved, issues could be investigated quickly, alternatives were considered, and decisions could move into execution.

This experience taught me that the value of agents does not come primarily from the intelligence of the model. It comes from the quality of the operating environment built around it. Clear objectives, complete context, reliable records, repeatable standards, and disciplined human review matter as much as raw model capability.

It also taught me that agents expose organizational weakness. If policies conflict, records are scattered, ownership is unclear, or leaders have not defined success, the agent will not magically resolve those conditions. It may reproduce them faster. Building agents therefore becomes an exercise in organizational clarity.

How I Am Applying Agents Across My Work

Educational leadership and student outcomes

Within a technology career school, an agent-powered approach can connect strategy to the details that determine student outcomes. Specialized agents can help reconcile recruitment records, analyze why prospects do not complete an application, prepare outreach sequences, track operational work, review program materials, monitor deadlines, organize evidence, and support staff with consistent information.

The objective is not automated education without teachers or counselors. It is to remove avoidable administrative friction so people can spend more time on judgment, instruction, encouragement, and relationships. An outreach agent can identify follow-up gaps, but a person should handle sensitive conversations. A policy agent can locate the applicable requirement, but institutional leadership remains accountable for interpretation and compliance. A curriculum agent can compare materials with certification objectives, but qualified educators must approve what students are taught.

Business strategy and consulting

In technology consulting and business development, agents help convert incomplete requests into structured decisions. They can organize discovery notes, identify unanswered questions, classify an engagement, estimate provisional effort from approved benchmarks, prepare

solution outlines, flag dependencies, and maintain continuity as an opportunity evolves. The result is a faster and more disciplined path from ambiguity to a reviewable proposal.

Confidentiality is essential. Client-specific information must remain within approved systems and access boundaries. Agents should not make commercial commitments, approve security designs, or replace accountable engineering review. Their role is to increase the quality and speed of preparation while leaving authorization with the responsible professionals.

Thought leadership and public communication

Agents also expand my capacity to develop ideas and bring them to the public. Research, drafting, editing, source validation, visual planning, website publishing, campaign preparation, and performance analysis are separate disciplines. A coordinated set of agents can support each stage while preserving a coherent thesis and voice.

The first obligation is authenticity. An agent may help articulate an idea, but it should not manufacture experience or conviction. I remain responsible for the claims I make, the values I express, and the advice attached to my name. The agent accelerates the work; it does not become the authorial conscience.

Personal objectives and executive capacity

The same principles apply outside a formal organization. A personal agent portfolio can support career decisions, financial planning, learning, health organization, family commitments, writing, major purchases, and long-term goals. These agents can preserve context, compare alternatives, prepare questions for professionals, and help turn intention into a sequence of actions.

Personal agents require strong boundaries. They should not be treated as physicians, attorneys, financial fiduciaries, pastors, or substitutes for trusted relationships. Their appropriate role is to help a person prepare, organize, reflect, and follow through while qualified humans remain responsible for high-stakes professional judgments.

The AGENT Portfolio Framework

The AGENT Portfolio Framework provides a leadership method for creating an agent that serves a real objective and remains governable. Every agent should be defined through five elements before it is given meaningful authority.

Element	Core question	Required design decision
A Aim	What strategic outcome must this	Define the objective users

	agent advance?	deliverables boundaries and value hypothesis.
G Guardrails	What may it do and what requires approval?	Set data rules permissions escalation thresholds prohibited actions and human checkpoints.
E Expertise	What must it know and what tools does it need?	Connect trusted knowledge instructions models systems and role-specific methods.
N Network	With whom and what must it coordinate?	Design agent handoffs human ownership system interfaces and exception paths.
T Tracking	How will we know whether it performs well?	Establish tests outcome measures quality thresholds logs reviews and improvement cycles.

Aim

Begin with an outcome, not a fascination with the technology. “Build an agent” is not an aim. “Reduce the time required to produce a complete, reviewable opportunity brief while improving consistency” is an aim. The better the aim, the easier it becomes to define scope and measure value.

Every aim should identify the user, the decision or deliverable, the desired improvement, and what remains outside the agent’s responsibility. A narrow, valuable first use case is usually better than an ambitious agent expected to transform an entire department.

Guardrails

Guardrails translate responsibility into system design. They determine what information the agent may access, which tools it may use, what actions it may take, what content it must refuse, when it must ask for help, and which results require human approval.

OpenAI’s practical guidance recommends layered guardrails rather than relying on a single control. NIST’s work on agent standards, identity, authorization, and AI risk management reinforces the same principle: trustworthy agents require secure identities, limited permissions, monitoring, and governance appropriate to their real-world impact.

Expertise

An agent's expertise is constructed. It includes the model, but also the instructions, examples, institutional vocabulary, policies, records, tools, and methods available to it. General intelligence without current organizational context produces generic work. Context without reliable source control produces confidently inconsistent work.

Expertise must therefore have provenance. Leaders should know where the agent obtained a fact, whether the source is current, and which information takes precedence when records conflict. An agent that cannot distinguish policy from opinion should not be entrusted with policy decisions.

Network

No meaningful agent operates alone. It interacts with people, data sources, business systems, and sometimes other agents. Network design specifies the handoffs. A research agent may send evidence to a strategy agent. The strategy agent may send an approved recommendation to an operations agent. The operations agent may create tasks but require a human to authorize external communication.

The network should not become complex merely because multi-agent technology is available. OpenAI's agent-building guidance advises teams to maximize a single agent's capabilities before adding complexity. Multiple agents are justified when specialization, separation of duties, independent review, or tool boundaries create clear value.

Tracking

An agent is not ready because it produced one impressive demonstration. It is ready when performance can be observed across representative work. Tracking should include output quality, task completion, time, cost, exception rate, escalation behavior, user trust, and business outcomes.

Evaluations should test both normal and adversarial conditions. NIST has highlighted agent hijacking and the possibility that agents may be manipulated through untrusted content. Testing must therefore include misleading instructions, missing data, conflicting sources, attempted permission escalation, and situations in which the correct behavior is to stop and ask for help.

Designing Agents That Perform Well

Reliable agents are designed around work, not demonstrations. The following principles have proven most important in my own use and align with emerging industry guidance.

Give the agent a job description

State the mission, users, recurring responsibilities, sources of truth, expected outputs, prohibited actions, escalation conditions, and definition of done. A useful agent role resembles a strong job description combined with an operating procedure.

Build context deliberately

Persistent context is a strategic asset. Preserve terminology, decisions, preferences, templates, and historical facts that the agent needs. Remove obsolete material and distinguish current policy from archived information. Context should improve continuity without becoming an uncontrolled warehouse of sensitive data.

Design structured handoffs

When work moves between agents or people, specify what must accompany it: the objective, current status, evidence, assumptions, unresolved questions, risks, and requested decision. Weak handoffs cause the next participant to reconstruct the problem and often lose the reasoning that produced the recommendation.

Require evidence where evidence matters

Research and compliance agents should cite authoritative sources. Analytical agents should expose assumptions and calculations. Operational agents should identify the record that supports a status. The appropriate standard varies by task, but important decisions should not rest on unverifiable assertions.

Use human approval according to consequence

Human review should be risk-based. A low-risk internal summary may require only periodic sampling. A public statement, payment, employment decision, regulatory submission, security change, or message to a vulnerable person should require explicit approval. The purpose of human review is not ceremonial oversight; it is accountable judgment at the point of consequence.

Evaluate the system not just the model

Failures may originate in the model, instructions, source material, tool configuration, permissions, handoff logic, or human process. Replacing the model will not repair an unclear policy or broken data flow. Evaluation should isolate where the system failed and what change would prevent recurrence.

The Human Operating System

As agents become more capable, the question is not whether humans remain involved. The question is where human contribution becomes most valuable. In an agent-powered organization, people move upward from repetitive production toward purpose, judgment, relationships, exception handling, and system improvement.

Five responsibilities must remain unmistakably human and accountable:

- **Vision:** deciding what future is worth pursuing and whose interests the system should serve.
- **Moral judgment:** resolving questions that cannot be reduced to efficiency or probability.
- **Accountability:** owning decisions and consequences rather than attributing them to an algorithm.
- **Relationship:** earning trust, showing empathy, negotiating meaning, and leading people through change.
- **Stewardship:** ensuring that short-term productivity does not undermine dignity, security, fairness, or long-term capacity.

This is especially important for faith-based and mission-driven leaders. Intelligence and efficiency are not the highest standards by which work should be judged. Technology should remain subject to truth, service, dignity, responsibility, and concern for people who may bear the cost of change without sharing its benefits.

The agent-powered enterprise needs stronger human leadership, not less. When digital systems can act more quickly and across greater scale, unclear values and weak governance also scale. Leadership supplies the direction that intelligence alone cannot.

Governance Before Autonomy

Organizations should not grant autonomy faster than they build governance. The risk profile changes when an AI system moves from generating text to accessing records, sending messages, changing configurations, approving transactions, or coordinating other agents.

A minimum governance model should address:

Governance domain	Leadership requirement
Ownership	Assign a named business owner and technical custodian for each production agent.
Identity	Give agents distinct identities where possible rather than shared human credentials.

Least privilege	Grant only the data and actions required for the defined role.
Data protection	Classify information and prohibit unapproved use of confidential regulated or personal data.
Human approval	Define consequential actions that cannot proceed without authorization.
Logging	Record important inputs tool calls decisions outputs exceptions and approvals.
Evaluation	Test quality security and policy compliance before deployment and after changes.
Incident response	Provide a way to suspend the agent investigate failures and correct affected work.
Lifecycle	Review agents periodically and retire duplicate obsolete or unjustified capabilities.

Agent sprawl is a foreseeable management problem. When departments build overlapping agents without shared standards, organizations inherit inconsistent instructions, duplicated costs, conflicting outputs, excessive permissions, and no clear owner. An agent registry should identify every operational agent, its aim, owner, knowledge sources, tools, permissions, evaluation status, and last review date.

Security testing must recognize that external documents, websites, emails, and tool outputs may contain instructions designed to manipulate an agent. Agents should treat untrusted content as data, not authority. High-impact systems require isolation, allowlisted actions, confirmation steps, and monitoring proportionate to the possible harm.

Measuring the AI Workforce

The most persuasive evidence for an agent is not that people enjoy using it. It is that the agent improves an outcome while remaining within acceptable risk and cost. Measurement should begin before deployment so the prior process can be compared with the new one.

Dimension	Example measure	Why it matters
Quality	Rubric score factual accuracy defect rate reviewer acceptance	Prevents speed from disguising poor work.
Completion	Percentage of tasks completed without avoidable failure	Shows whether the agent reliably reaches the objective.

Efficiency	Cycle time human minutes cost per completed task	Quantifies capacity and economic value.
Escalation	Frequency reason and appropriateness of handoffs	Reveals uncertainty policy gaps and overreach.
Business impact	Enrollment revenue service response risk reduction or milestone attainment	Connects activity to strategy.
Trust and adoption	Use rate override rate satisfaction and reported incidents	Shows whether people can depend on the system.
Governance	Permission exceptions unsupported claims privacy events and audit findings	Makes responsible operation measurable.

Metrics must be interpreted together. A lower escalation rate is not positive if the agent is acting beyond its competence. A faster response is not valuable if the answer is wrong. High adoption may increase risk if users trust outputs uncritically. The scorecard should reward appropriate restraint as well as successful action.

A 90 Day Agent Portfolio Plan

Organizations do not need to begin with an enterprise-wide autonomous system. They need a disciplined sequence that produces evidence, develops leadership capability, and protects trust.

Period	Leadership focus	Key outputs
Days 1 to 15	Discover and prioritize	Work inventory pain points risk screen baseline measures and one selected use case.
Days 16 to 30	Design with AGENT	Agent charter guardrails knowledge map tool plan handoffs and evaluation set.
Days 31 to 45	Build a supervised pilot	Working agent limited permissions representative test cases and logging.
Days 46 to 60	Evaluate and red team	Quality results adversarial tests failure analysis user feedback and revisions.
Days 61 to 75	Deploy with controls	Named owner training approval

		workflow monitoring and incident procedure.
Days 76 to 90	Measure and decide	Outcome comparison ROI assessment governance review and scale revise or retire decision.

Days 1 to 15 Discover and prioritize

Inventory repetitive, document-heavy, research-intensive, or coordination-heavy work. Look for tasks with high volume, long delays, inconsistent quality, or excessive context switching. Avoid beginning with the most consequential decision in the organization. Select a use case that matters but can be safely supervised.

Days 16 to 30 Design with AGENT

Complete an AGENT charter. Establish the aim, guardrails, expertise, network, and tracking plan. Identify current authoritative sources and resolve conflicts before loading them into the system. Decide which actions remain read-only, which may be drafted, and which require approval.

Days 31 to 45 Build a supervised pilot

Create the smallest working agent that can address the use case. Begin with limited tools and permissions. Run it on historical or simulated work before live deployment. Capture every failure, correction, and missing piece of context.

Days 46 to 60 Evaluate and red team

Test normal cases, edge cases, incomplete information, conflicting records, malicious instructions, and requests outside scope. Include subject-matter experts who did not build the agent. The objective is not to prove the pilot works; it is to discover the conditions under which it does not.

Days 61 to 75 Deploy with controls

Introduce the agent to a limited group. Train users on its purpose, limitations, review obligations, and escalation path. Monitor both the agent and the people using it. Overtrust, workarounds, and unreported corrections are operational risks.

Days 76 to 90 Measure and decide

Compare results with the baseline. Determine whether the agent improved quality, capacity, speed, or outcomes enough to justify ongoing cost and risk. Expand permissions or scale only when supported by evidence. Revise agents that show a credible path to value. Retire those that do not.

The Agent Powered Individual

The agent portfolio is not only an enterprise concept. An individual can create a personal operating system in which specialized agents help manage the complexity of modern life. The key is to organize agents around enduring objectives rather than creating a new conversation for every task.

Personal agent	Purpose	Human boundary
Chief of staff	Organize priorities commitments decisions and follow up	The individual decides what matters and makes commitments.
Career strategist	Evaluate opportunities strengthen positioning and prepare interviews	The individual owns career choices and represents experience honestly.
Learning coach	Create plans explain concepts test understanding and track progress	The learner completes the work and verifies consequential knowledge.
Financial organizer	Categorize information model scenarios and prepare professional questions	Licensed professionals handle regulated advice where required.
Health organizer	Track questions symptoms appointments and clinician instructions	Medical diagnosis and treatment remain with qualified clinicians.
Writing and publishing partner	Research outline draft edit and adapt material	The author approves claims sources and final voice.
Personal reflection partner	Clarify options values tensions and next steps	AI does not replace faith community family friendship or counseling.

A personal portfolio should include privacy rules, source checks, and periodic cleanup. Convenience can lead people to provide more sensitive information than necessary. Personal agents should receive only the information required for the task, and important records should remain in systems designed to protect them.

What Leaders Should Do Now

The agent era is arriving before most organizations have established the vocabulary, governance, or management capability required for it. Waiting for the technology to become settled is not a strategy; neither is racing into autonomy without controls. Leaders should take five actions now.

- Develop personal fluency by using agents on meaningful but supervised work.
- Identify one workflow where an agent can improve a measurable outcome within acceptable risk.
- Establish an agent registry and minimum governance standard before proliferation begins.
- Train managers to design work, set quality standards, supervise exceptions, and evaluate agent performance.
- Prepare employees to move toward judgment, relationships, creativity, oversight, and continuous improvement.

Education and workforce systems must also respond. Future professionals will need more than AI awareness. They must learn process analysis, data literacy, tool integration, evaluation, cybersecurity, governance, and human-centered leadership. The ability to manage a collection of agents may become as fundamental to professional work as the ability to manage projects, information, and teams.

Conclusion The Future Belongs to Agent Orchestrators

Artificial intelligence is creating a new layer of organizational capacity. Agents can research, reason, coordinate, create, monitor, and act across workflows that once depended entirely on scarce human attention. But capacity without direction does not become progress. It becomes noise, duplication, and risk.

The defining capability of the next era will be orchestration: the ability to organize human talent and digital intelligence around worthwhile objectives. Agent orchestrators will know what to delegate, what to automate, what to reserve for people, and where to require accountability. They will treat context as infrastructure, guardrails as leadership, evaluation as management, and human dignity as a design requirement.

My own experience has convinced me that an agent portfolio can materially expand what one leader and one organization can accomplish. It can preserve institutional memory, accelerate decisions, strengthen execution, and make ambitious objectives more manageable. It can also reveal gaps in policy, data, ownership, and judgment that technology alone cannot repair.

The future of work will not belong simply to people who use AI. It will belong to those who can build and manage high-performing systems of people and agents—and who possess the wisdom to ensure that those systems serve human beings rather than diminish them.

The opportunity before us is not to remove humanity from work. It is to use intelligence, in all its forms, to increase human capacity, extend responsible leadership, and create a future in which more people can contribute, adapt, and thrive.

Appendix A Agent Charter

AGENT element	Questions to answer
Aim	What outcome will the agent advance? Who uses it? What is in and out of scope? What is the baseline?
Guardrails	What data may it access? What actions may it take? What requires approval? When must it stop or escalate?
Expertise	What instructions sources examples models and tools does it need? Which sources control when information conflicts?
Network	Who owns it? What people agents and systems are involved? What information must accompany each handoff?
Tracking	How will quality completion cost speed risk adoption and business impact be evaluated?

Agent name _____

Business owner _____

Technical custodian _____

Review frequency _____

Current authority level Draft only Supervised action Limited autonomy

Appendix B Agent Performance Scorecard

Measure	Baseline	Target	Current	Evidence and action
Quality score				
Completion rate				
Cycle time				
Human review time				
Cost per completed task				
Escalation rate				
Unsupported claim rate				
Permission or privacy incidents				
User adoption				
Business outcome				

Selected Sources

1. OpenAI. A Practical Guide to Building AI Agents. OpenAI Business, 2025.
<https://openai.com/business/guides-and-resources/a-practical-guide-to-building-ai-agents/>
2. OpenAI. Agents SDK and agent development guidance. OpenAI Developers.
<https://developers.openai.com/api/docs/guides/agents>
3. OpenAI. New Tools for Building Agents. March 11, 2025. <https://openai.com/index/new-tools-for-building-agents/>
4. Microsoft. The 2025 Annual Work Trend Index The Frontier Firm Is Born. April 23, 2025.
<https://www.microsoft.com/en-us/worklab/work-trend-index/2025-the-year-the-frontier-firm-is-born>

5. Microsoft. 2026 Work Trend Index Agents Human Agency and the Opportunity for Every Organization. May 5, 2026.
<https://www.microsoft.com/en-us/worklab/work-trend-index/agents-human-agency-and-the-opportunity-for-every-organization>
6. National Institute of Standards and Technology. AI Risk Management Framework.
<https://www.nist.gov/itl/ai-risk-management-framework>
7. National Institute of Standards and Technology. AI Agent Standards Initiative. February 17, 2026. <https://www.nist.gov/artificial-intelligence/ai-agent-standards-initiative>
8. National Cybersecurity Center of Excellence. Software and AI Agent Identity and Authorization. <https://www.nccoe.nist.gov/projects/software-and-ai-agent-identity-and-authorization>
9. NIST Center for AI Standards and Innovation. Strengthening AI Agent Hijacking Evaluations. January 17, 2025. <https://www.nist.gov/news-events/news/2025/01/technical-blog-strengthening-ai-agent-hijacking-evaluations>
10. Anthropic. How AI Is Transforming Work at Anthropic. 2025.
<https://www.anthropic.com/research/how-ai-is-transforming-work-at-anthropic>

About the Author

Therelee D. Washington II is a technology leader, entrepreneur, educator, consultant, and AI strategist with more than two decades of experience helping organizations apply technology to business and workforce challenges. He is the founder of Texas Premier Technology Institute and AgiliShare Solutions Group. His work focuses on artificial intelligence, enterprise technology, workforce development, education, leadership, and the responsible use of innovation to expand human opportunity.

His prior white paper, *The Human Future in an Age of Dominant AI*, presented a practical roadmap for careers, employers, educators, and public leaders navigating AI-driven change. In this paper, he extends that work from preparation to execution by examining how leaders can build and govern agent-powered operating models.