



Measuring Modern AI Laptops: Real Productivity and Battery Life for the Road Warrior

The Integrated NPU in the AMD Ryzen AI 7 PRO 450
Brings All-Day AI to the Mobile Professional

Ryan Shrout

In partnership with:



signal65.com

Contents

3	How Road Warriors Benefit from AI-Capable Laptops	17	Click-to-Do Slide Update
4	The Road Warrior	19	Research Client Documents
6	Testing Configuration	21	Generate Tasks from Emails
7	The NPU Efficiency Story	23	Implications for Road Warriors with AI Laptops
9	Compiling the Real-World Measurements	25	Appendix: Our Testing Theory
11	The Silicon Generation Story	27	Important Information About this Report
13	Document Draft	28	System Configurations
15	Create Meeting Notes		

How Road Warriors Benefit from AI-Capable Laptops



The laptop is the road warrior's entire office, and in 2026 the AI-capable laptop changes what that office can do away from a desk and away from a charger. Modern AMD Ryzen AI laptops run AI models directly on the same silicon that powers everyday productivity, so the work does not wait on a cloud round-trip or a reliable connection. For professionals who spend the bulk of their day across email, document creation, meetings, and client research, the question is no longer whether AI tools will reach the laptop, but how much value the right hardware unlocks.

Building on our previous analysis of office workers and road warriors we have done at Signal65, this study answers two questions at once. How much time does AI save the mobile professional, and what does the right silicon do for both speed and battery life under real multitasking? To answer them, we ran five common road warrior workflows on two AMD laptop systems, a current-generation Ryzen AI 7 PRO 450 laptop and a prior-generation Ryzen 5 PRO 6650U laptop. Each workflow was tested with AI assistance on both systems and with a fully manual approach. We then measured productivity and battery use under

concurrent multitasking, comparing the NPU and GPU execution paths on the new machine and the new machine against the old.

The results show that where the AI runs matters as much as how fast it runs. Under concurrent multitasking, moving the local model from the GPU to the integrated NPU raised the office-productivity score by 11% while drawing 50% less battery capacity, and against the prior-generation laptop the new machine delivered 1.8x the productivity on just over a third as much battery. The time savings are substantial as well. Across all five workflows, the new Ryzen AI PRO laptop completed in 17.4 minutes a representative bundle of common road warrior tasks that took 62.8 minutes to do manually, a 72% reduction, and it ran the same AI workloads 1.9x faster than the prior-generation laptop.

The case for current-generation AI laptops is both immediate and broader than it appears. Immediate, because the productivity and battery gains shown here are available now on shipping hardware and shipping software. Broader, because the same silicon trends that enable today's workflows are accelerating, and software increasingly assumes capable on-device compute and a modern memory

subsystem to deliver AI experiences. Mobile professionals running prior-generation laptops will increasingly find themselves limited not by their willingness to adopt AI tools but by the hardware those tools require.

Key Highlights:

 Running AI on the integrated NPU delivers **11% higher productivity on 50% less battery.**

 The new AMD Ryzen AI laptop **reduces task time by 72%** across our road warrior workflows.

 Workflow steps like meeting summarization run up to **17.8x as fast** with AI than manual work.

 The new Ryzen AI CPU runs the **same AI workload 1.9x as fast** as the prior-generation AMD laptop while **drawing just a third of the battery** capacity.

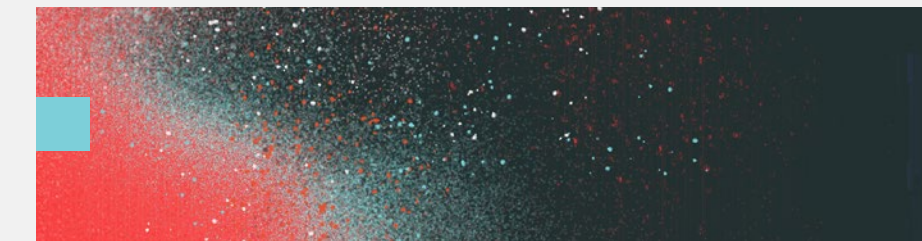
The Road Warrior

The Road Warrior persona covers mobile professionals who work primarily from a laptop, often away from a desk, frequently on battery, and across the same integrated suite of business applications as the office worker but under tighter power and connectivity constraints. This persona lives on the productivity of three things at once: the laptop hardware, the speed of the AI tools that laptop can run, and the battery available to run them. The persona encompasses a range of different kinds of workers:

- Field sales representatives and sales engineers who work between customer sites, airports, and hotels
- Consultants and client-facing analysts who review client material, draft deliverables, and summarize meetings on the move

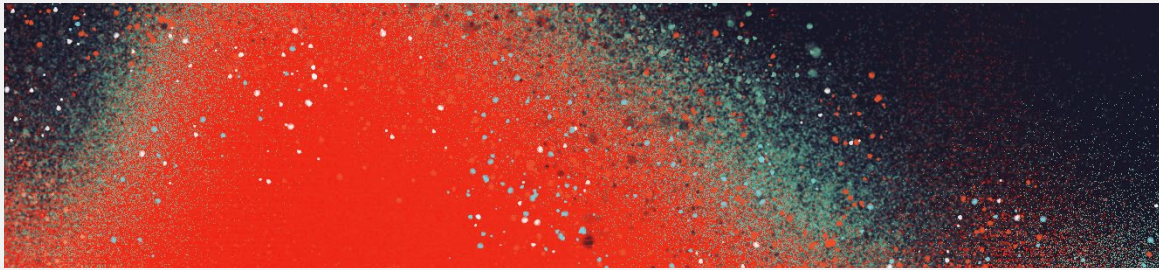
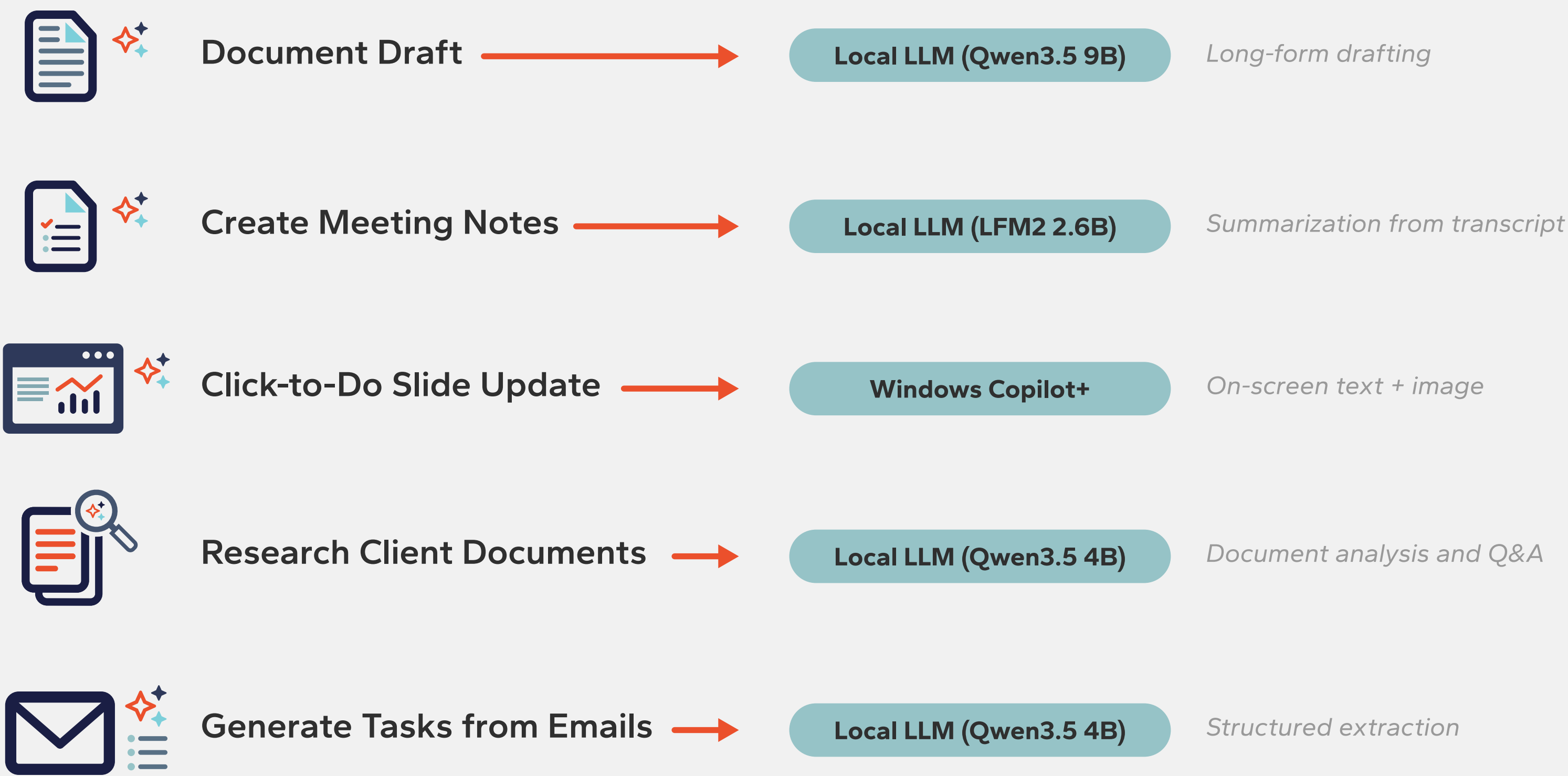
- Executives who travel and need to stay productive across flights, lobbies, and back-to-back meetings
- Any knowledge worker whose most productive hours happen away from reliable power and connectivity
- Workers whose individual output compounds directly into team and organizational throughput, wherever they happen to be working

What separates this persona from the desktop office worker we evaluated earlier this year are a set of constraints. The road warrior contends with limited battery, intermittent power, variable connectivity, paired with the frequent need to multitask, running a video call while doing live AI work, without the machine bogging down or dying. The workflow spans content drafting (Document Draft), summarization (Create Meeting Notes), multi-application UI interactions (Click-to-Do Slide Update), document analysis (Research Client Documents), and structured extraction (Generate Tasks from Emails). Four of the five workflows use local large language models running on the laptop. One uses the Windows Copilot+ Click-to-Do feature. Keeping inference on-device matters even more for the mobile worker, who may be working without reliable connectivity.



REAL PRODUCTIVITY AND ALL-DAY BATTERY
FOR THE ROAD WARRIOR

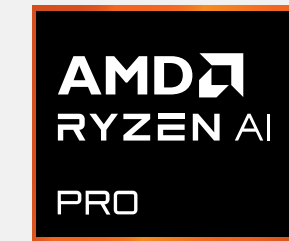
The Road Warrior



Testing Configuration

Our testing uses two different AMD laptop systems to look at two primary comparisons. The first is the productivity gain from AI assistance over manual workflows. The second is the gain from current-generation silicon over prior-generation silicon on the same AI-based workload. The latest generation system is an HP EliteBook X G2a built on the AMD Ryzen AI 7 PRO 450, which integrates a dedicated NPU alongside the Ryzen CPU and Radeon GPU. The reference prior-generation system is an HP EliteBook 845 G9 built on the AMD Ryzen 5 PRO 6650U, a capable mobile CPU from a prior platform generation that predates the integration of AI-specific silicon. Both systems were configured with 32GB of RAM, with the new system using LPDDR5X memory and the older

system using DDR5. The manual workflow is hardware-independent and applies equally to either system in our experience, providing a stable baseline against which both AI configurations can be compared.



	MODERN SYSTEM - HP ELITEBOOK X G2A	PRIOR-GEN REFERENCE - HP ELITEBOOK 845 G9
CPU	AMD Ryzen AI 7 PRO 450	AMD Ryzen 5 PRO 6650U
Integrated NPU	AMD XDNA2	Not present
Graphics	AMD Radeon (integrated)	AMD Radeon (integrated)
RAM	32GB LPDDR5X	32GB DDR5
Storage	1TB NVMe SSD	1TB NVMe SSD
Operating System	Windows 11 Pro 26200	Windows 11 Pro 26200
Copilot+ PC Eligible	Yes	No

Each workflow was run three times on each applicable configuration and the results averaged. For local-LLM workflows, the same models, prompts, and inputs were used on both systems, with inference running on the integrated GPU through LM Studio on both machines so that the workflow comparison isolates the silicon generation rather than the execution path. The NPU execution path

is evaluated separately in the efficiency and battery testing described below. The Click-to-Do workflow was tested only on the new system because the prior-generation Ryzen 5 PRO 6650U does not meet the Copilot+ PC requirement for an integrated NPU at 40 TOPS or above.

This study also extends our methodology with a new dimension that a laptop makes possible and a desktop does not. Alongside the workflow timings, we measured productivity and battery use under concurrent multitasking, running an office-productivity benchmark at the same time as a local AI workload. Battery is reported as capacity drawn in milliwatt-hours, and we captured it separately for the NPU and GPU execution paths on the new machine and for the new machine against the old. This is what lets us show not just how fast the AI runs, but how much battery it costs to run it.

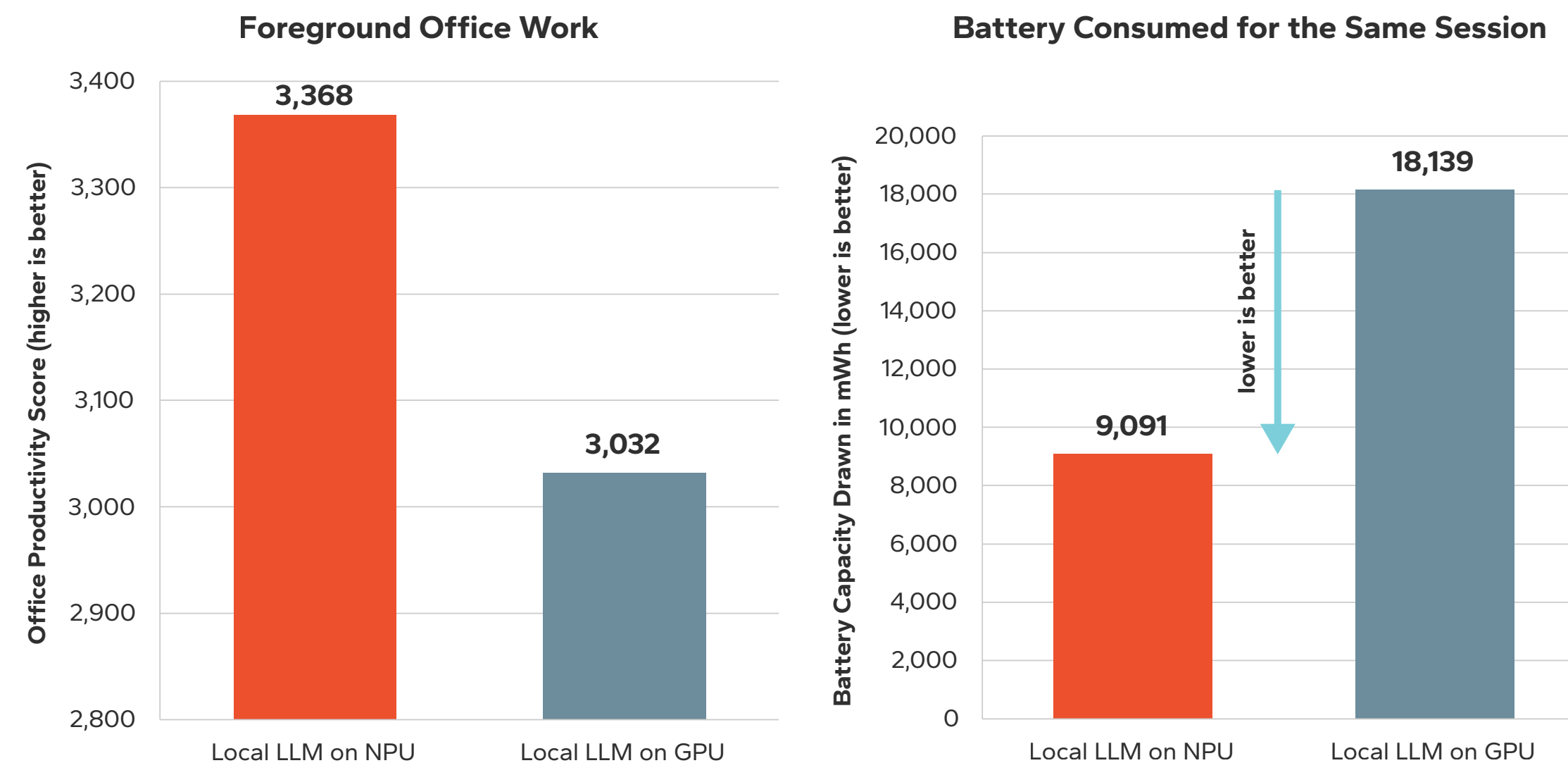
The NPU Efficiency Story

On a laptop, where the AI decides how long the work can continue

A road warrior's working day is bounded by battery as much as by hours. The workflow time savings measured later in this report only matter if the laptop is still running when the work needs doing, so the first question we put to the new hardware was not how fast it runs AI but what running AI costs in power. To measure that, we ran an office-productivity benchmark concurrently with a local AI workload, simulating the real multitasking a road warrior does, and recorded both the productivity score and the battery capacity drawn. The AI workload here is a sustained load rather than a timed task, so what we are measuring is what happens to foreground office work and to the battery while inference runs. On the same laptop running the same 9B model, moving that inference from the GPU to the

Where the AI Runs Matters - Same Laptop, Same Model (Qwen3.5 9B)

Office-productivity benchmark is the measured foreground workload; only the local LLM moves between the NPU and GPU.



integrated NPU raised the office-productivity score by 11% while drawing 50% less battery capacity. The foreground work got faster and the session cost roughly half as much power.

The advantage compounds when the new laptop is compared against the prior generation under the same multitasking load. Holding the execution path constant, with the local model on the GPU on both machines so that only the silicon differs, the new laptop delivered 1.6x the office-productivity score while drawing 25% less battery capacity. Comparing each system on the best path available to it, which means the NPU on the new laptop and the GPU on the prior generation, the new laptop delivered 1.8x the productivity score while drawing just over a third as much battery capacity, roughly 62% less. For the road warrior, this is the difference that matters. The new laptop does more work on a fraction of the battery, which can be the difference between arriving with a finished deliverable and a charged machine, or arriving with neither.

REAL PRODUCTIVITY AND ALL-DAY BATTERY
FOR THE ROAD WARRIOR

The NPU Efficiency Story

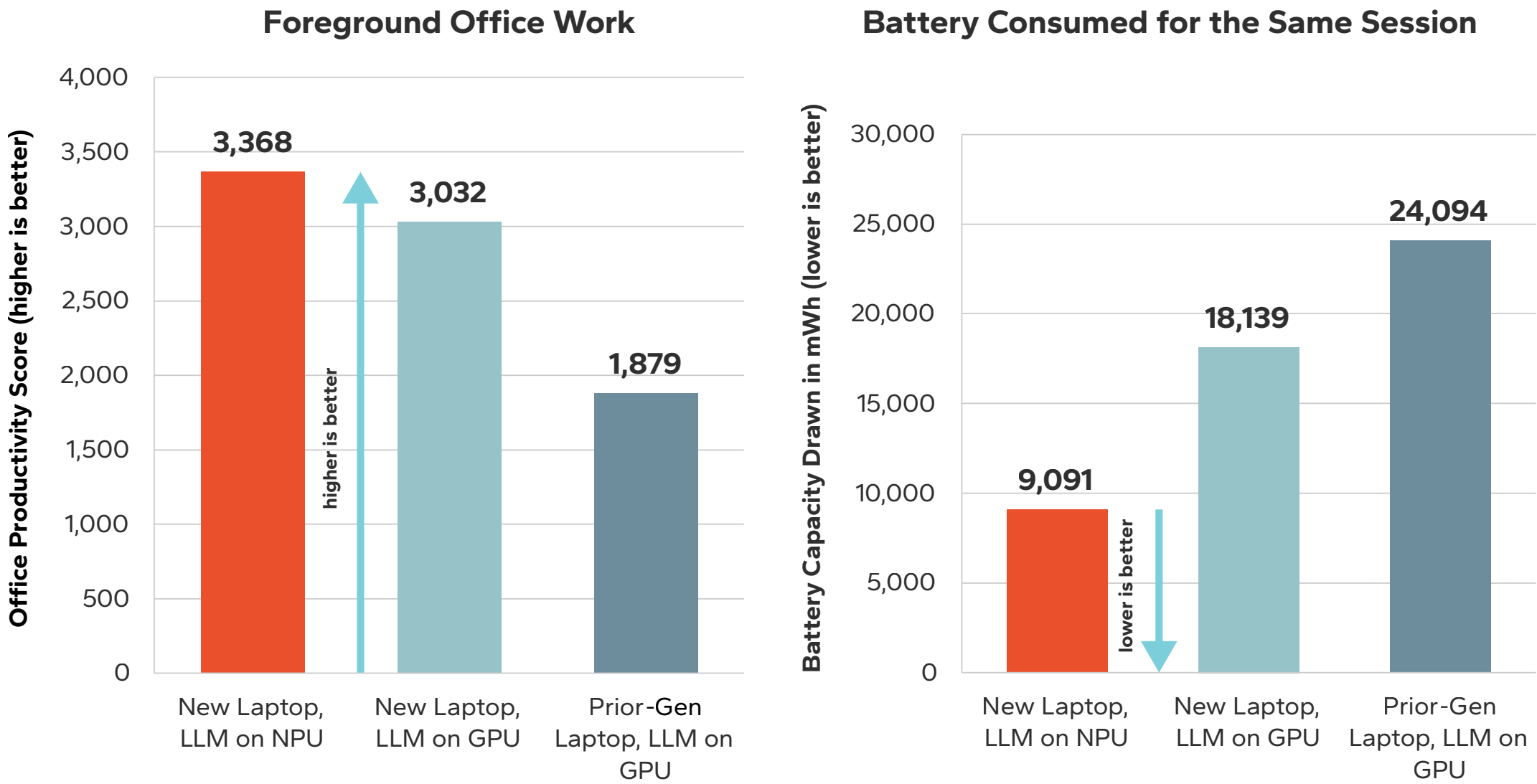
Comparison	New (Ryzen AI 7 PRO 450)	Reference	Result
Productivity score, same laptop, LLM on NPU vs GPU	3,368 (LLM on NPU)	3,032 (LLM on GPU)	+11% with LLM on NPU
Battery drawn, same laptop, LLM on NPU vs GPU	9,091 mWh (LLM on NPU)	18,139 mWh (LLM on GPU)	50% less with LLM on NPU
Productivity score, new vs prior-gen, both LLM on GPU	3,032 (LLM on GPU)	1,879 (LLM on GPU)	60% higher
Battery drawn, new vs prior-gen, both LLM on GPU	18,139 mWh (LLM on GPU)	24,094 mWh (LLM on GPU)	25% less
Productivity score, new vs prior-gen, best path each	3,368 (LLM on NPU)	1,879 (LLM on GPU)	80% higher
Battery drawn, new vs prior-gen, best path each	9,091 mWh (LLM on NPU)	24,094 mWh (LLM on GPU)	62% less

Refresh planning for mobile workers in 2026 and beyond should treat AI capability and on-battery efficiency as primary specifications rather than nice-to-haves. Specifying mobile systems without integrated NPUs and modern memory subsystems will leave both productivity and battery life on the table for any worker whose day includes AI-assisted steps.

What that efficiency buys is time on the work itself. The sections that follow measure five common road warrior workflows with and without AI assistance, and then isolate how much of the gain comes from the silicon generation.

New vs. Prior-Generation Laptop Under the Same AI Multitasking Load

Prior-generation laptop has no NPU, so its only execution path is the GPU; same Qwen3.5 9B model throughout.



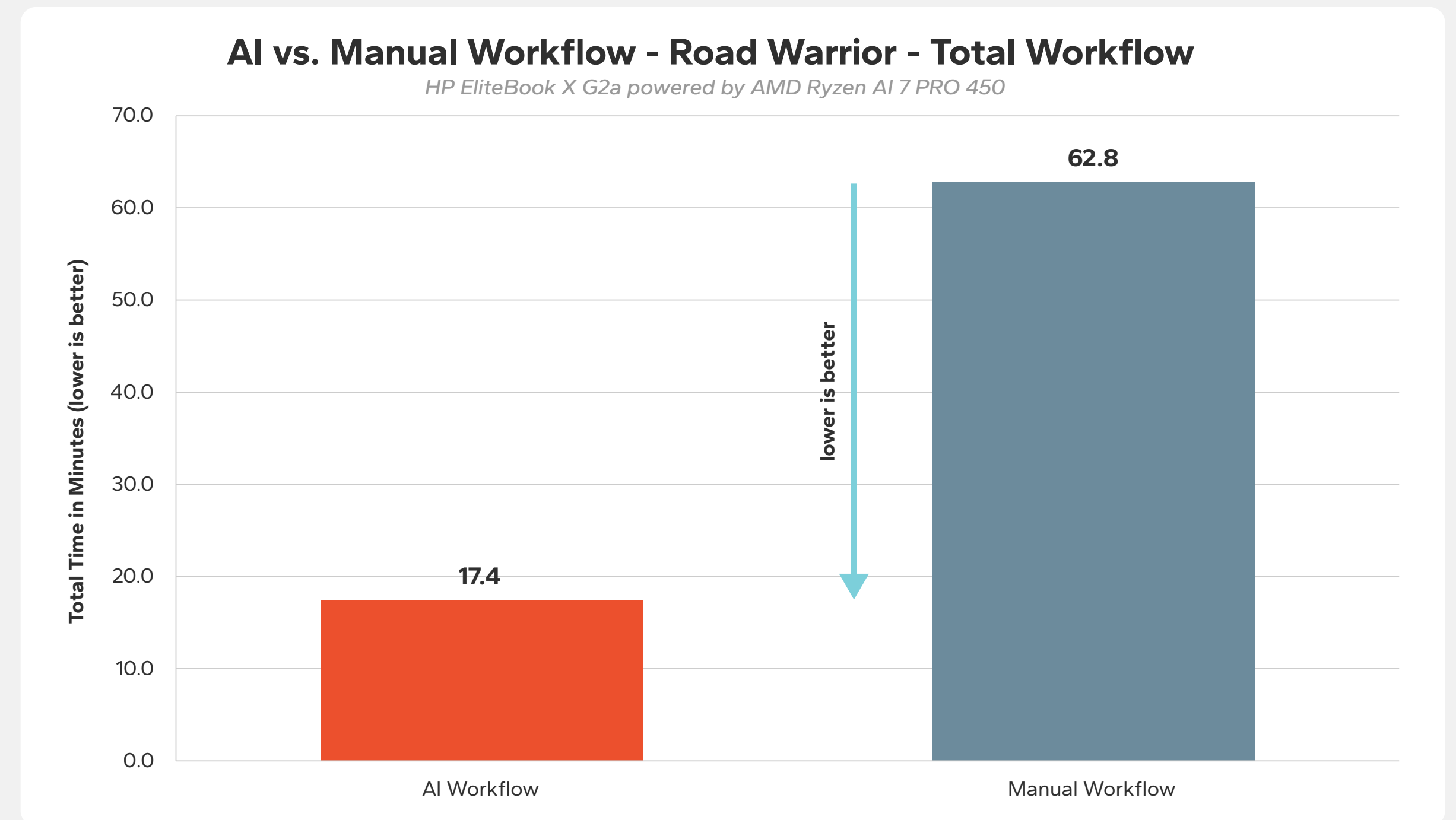
On the same laptop, moving the AI workload to the integrated NPU left more headroom for foreground office work and halved the battery usage of the session. Against the prior generation on each system’s best available path, the new laptop does **1.8x the work** on just over a third as much battery capacity.

Compiling the Real-World Measurements

Our testing shows that the new AMD laptop, paired with current AI tools, delivers measurable productivity gains across every workflow we measured. Across the complete Road Warrior test suite, AI-enabled workflows on the Ryzen AI 7 PRO 450 system completed in approximately 17 minutes, compared to nearly 63 minutes for the same outcomes produced manually. That is a 72% reduction in total task time on a representative bundle of common road warrior tasks, and a 3.6x speed-up across the five workflows. These gains are in no way marginal, and they are available today on shipping hardware and software.

Projecting these gains to a typical workweek illustrates the impact at scale. At realistic task frequencies (one document draft, three meetings, two slide updates, two research tasks, and three task batches

per day) the AI-enabled workflows save approximately 1.2 workdays per week. The savings are not evenly distributed across workflows. Meeting summarization alone accounts for more than two-thirds of the total weekly recovery because it combines a high task frequency with a workload that AI models are incredibly well-suited to handle. The other workflows show more modest but consistent gains in the 1.7x to 3.4x range, and those smaller gains still compound through daily repetition.



The Ryzen AI 7 PRO 450 system showed a **72% reduction in total task time**.

REAL PRODUCTIVITY AND ALL-DAY BATTERY
FOR THE ROAD WARRIOR

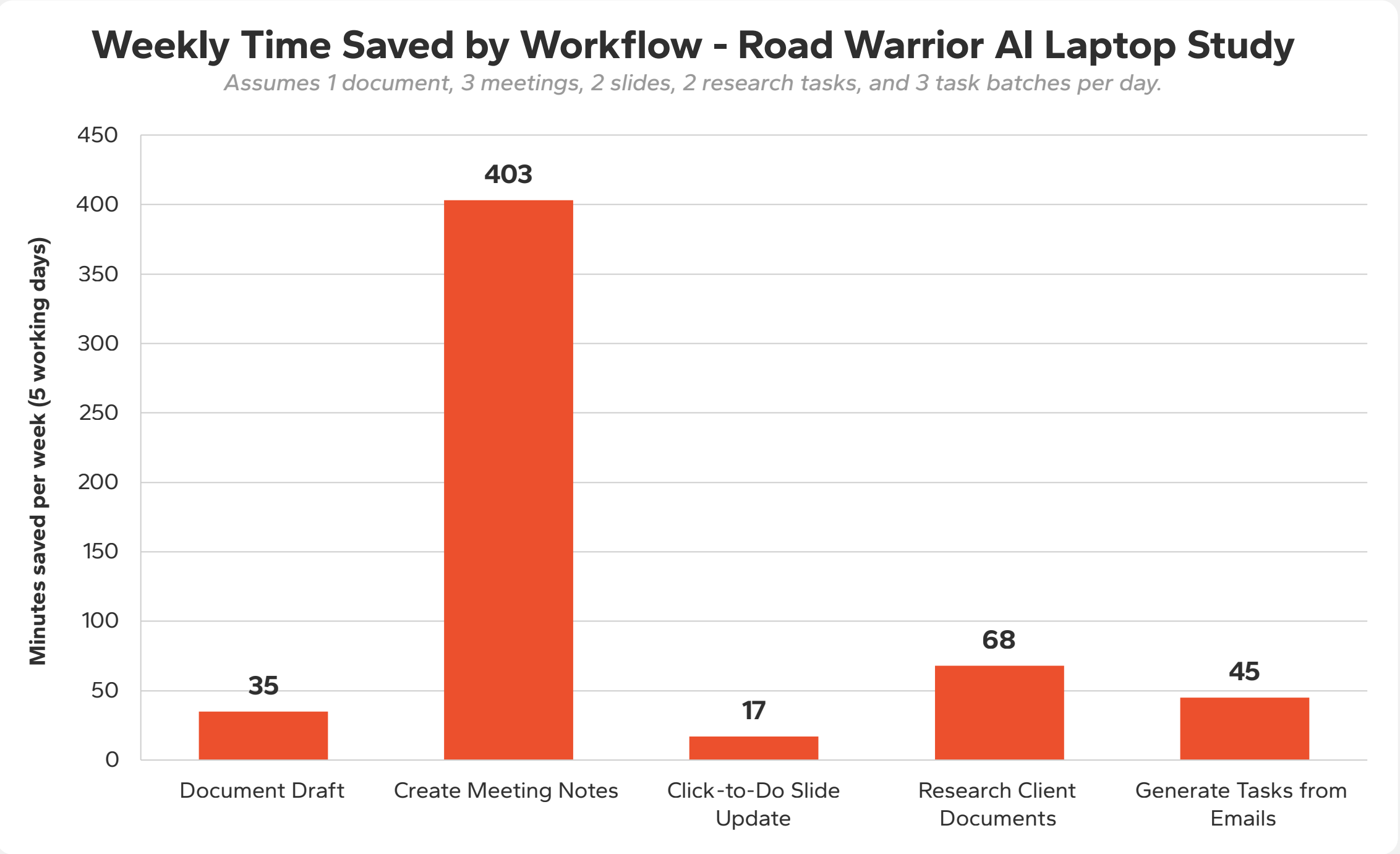
Compiling the Real-World Measurements

Workflow	Frequency	Per-Task Savings	Weekly Savings	Share of Total
Document Draft	1 / day	7.0 min	35 min	6%
Create Meeting Notes	3 / day	26.9 min	403 min	71%
Click-to-Do Slide Update	2 / day	1.7 min	17 min	3%
Research Client Documents	2 / day	6.8 min	68 min	12%
Generate Tasks from Emails	3 / day	3.0 min	45 min	8%
TOTAL			568 min / week	100%
Equivalent to approximately 9.5 hours per week, or 1.2 workdays per week recovered.				

The time savings are the easiest gains to measure, but the reduction in cognitive load may matter just as much for a professional working between meetings and across time zones. Summarization removes the sustained concentration required to work through a meeting transcript or a long client document. Structured extraction surfaces action items that would otherwise be missed in a crowded

inbox. Multi-application features such as Click-to-Do cut context switching, which carries its own well documented cost to knowledge work and is especially costly on a small mobile screen. A road warrior arriving at a client meeting has spent less of their finite attention getting there. The complete productivity story for AI-capable laptops is the sum of these effects, not just the raw seconds saved on each workflow.

Road warriors using AI tools on a modern AMD laptop **recover the equivalent of 1.2 workdays each week** without extending work hours or adding headcount, and they do it on battery.



The Silicon Generation Story

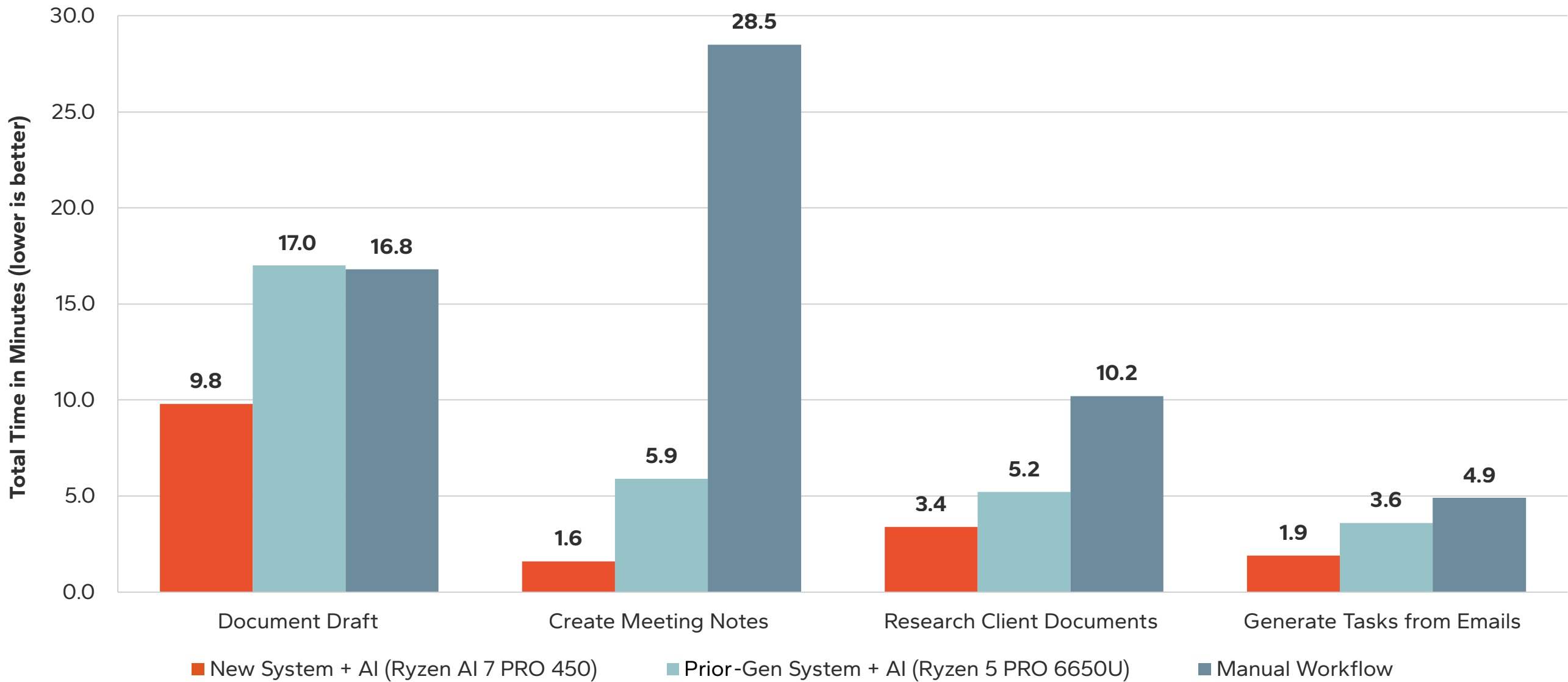
How much of the workflow gain comes from the hardware itself

The workflow measurements above show what AI assistance saves on the new laptop. They do not, on their own, separate the contribution of the AI tools from the contribution of the silicon running them.

To isolate the hardware contribution, we ran the same AI workflows on a prior-generation AMD Ryzen 5 PRO 6650U laptop and compared the results to the new Ryzen AI 7 PRO 450 laptop. The manual workflow stays identical on both systems and serves as a stable reference baseline. Across the four workflows tested on both systems, the new laptop ran the same AI workloads on average 90% faster than the prior-generation system, recovering 15 minutes per workflow set from hardware alone. New silicon does not just make AI possible, it makes AI dramatically better.

Hardware Contribution at Work - New vs. Prior-Gen Laptop on the Same AI Workflow

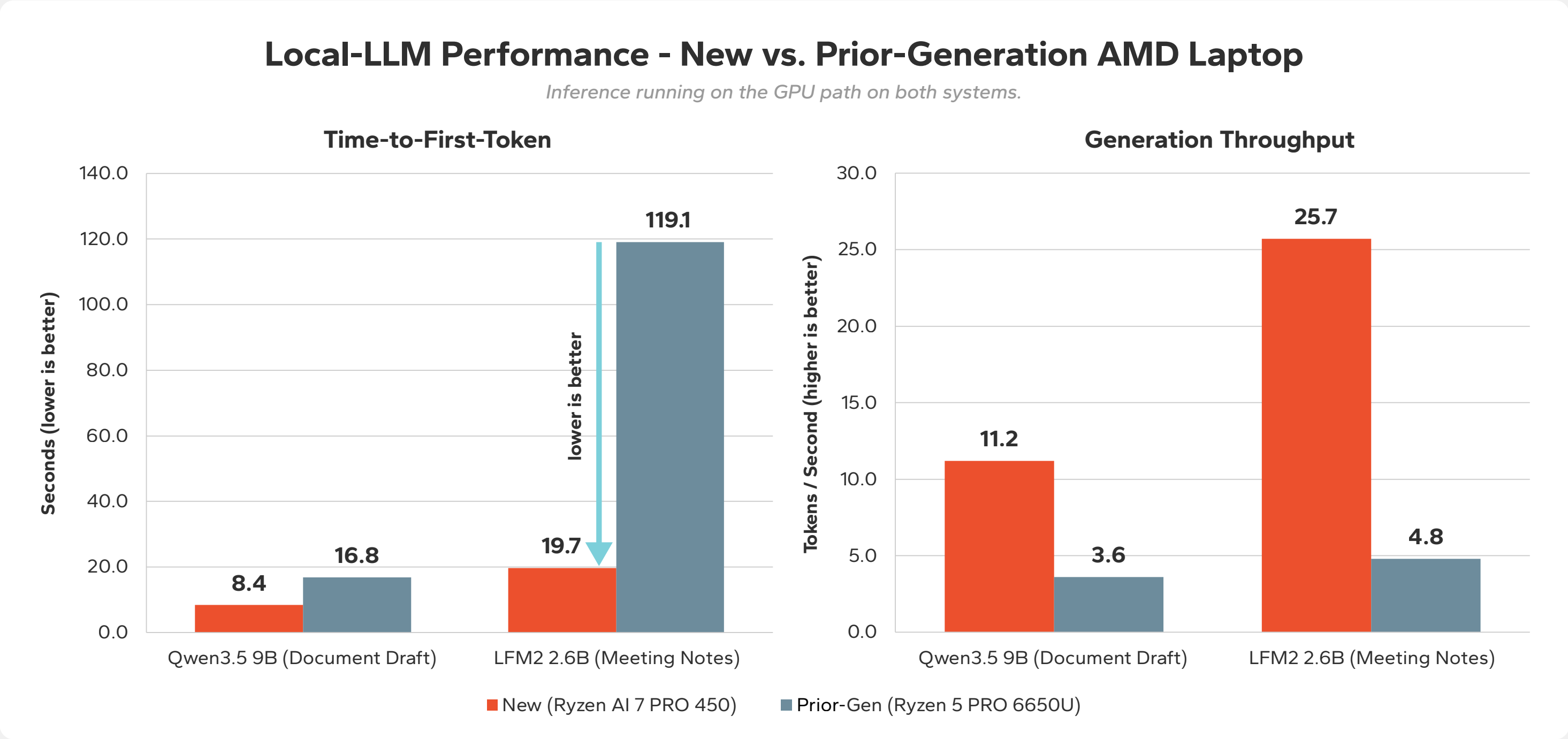
Local LLM inference runs on the GPU path on both systems; Click-to-Do excluded (not supported on prior-gen)



REAL PRODUCTIVITY AND ALL-DAY BATTERY
FOR THE ROAD WARRIOR

The Silicon Generation Story

On the local large language models used in this study, the new laptop delivers Time-to-First-Token between 2.0x and 6.0x faster than the older laptop, and generation throughput between 3.1x and 5.4x better. Time-to-First-Token determines how long the user waits before the model begins producing output, which is the most psychologically important latency for interactive AI work. Generation throughput determines how quickly the rest of the response arrives. For meeting summarization specifically, the gap is stark, with the new laptop reaching first token in 19.7 seconds versus 119.1 seconds on the old machine. Together these metrics explain why the new laptop widens the gap between AI and manual on every workflow.



Document Draft

Drafting a document such as a client proposal, a trip report, or a structured status memo is a common but important task for the mobile professional. The work involves taking source material (a brief, a set of notes, or structured bullet points) and producing a polished piece of writing. Our Document Draft scenario tests how much of that initial drafting can be productively offloaded to a local AI model, while keeping the manual polish step where human judgment is most valuable.

The AI-enabled workflow uses a 9B-parameter Qwen3.5 local model running on the AMD laptop through LM Studio. The user supplies a structured prompt with the source material and the model produces a complete draft. The human then performs an identical add-details step (adjusting voice, adding specifics, and final proofreading) on both the AI workflow and the manual workflow. The manual workflow simply skips the AI generation step and writes the draft from scratch before the same polish.

AI-Enabled Steps	Manual Steps
The AI version uses a local large language model to generate a complete first draft from a structured prompt, which the writer then polishes.	The manual version writes the entire draft from scratch and applies the same polish step at the end.
Open source material and target document	Open source material and target document
Launch local LLM (Qwen3.5 9B) via LM Studio	Draft the document from scratch in Word
Submit structured prompt with source material	
Receive and review AI-generated draft	
Apply manual polish (voice, specifics, proofreading)	Apply manual polish (voice, specifics, proofreading)

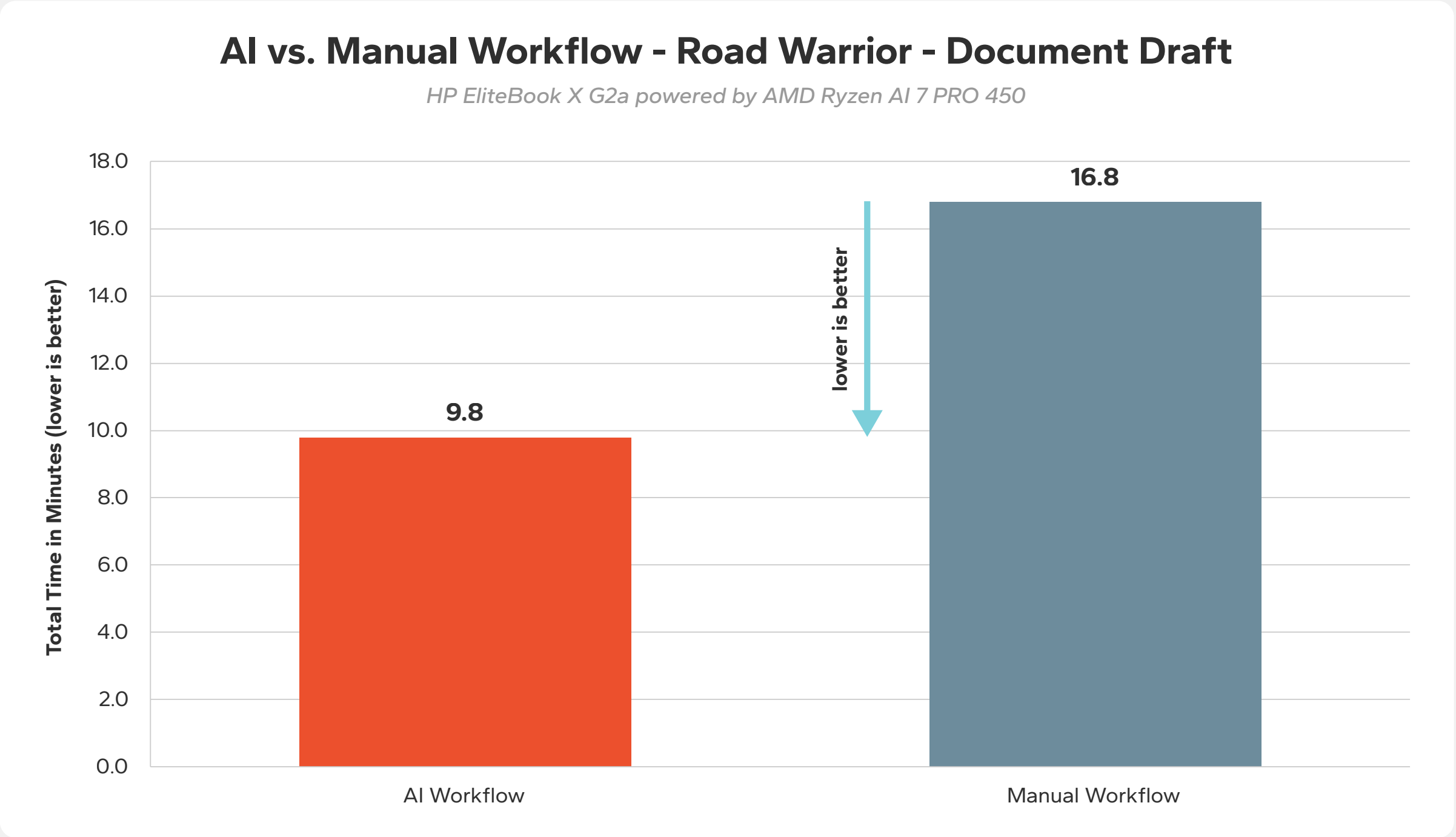
REAL PRODUCTIVITY AND ALL-DAY BATTERY
FOR THE ROAD WARRIOR

 Document Draft

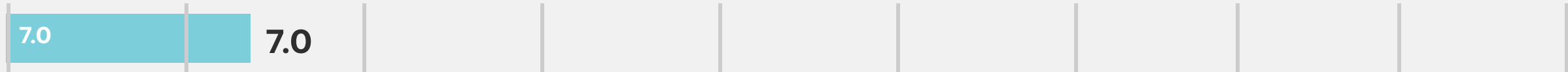
Performance results show that the AI-enabled workflow completes the entire document drafting process in 9.8 minutes, compared to 16.8 minutes for the manual approach, a 42% reduction in time and a 1.7x speed-up. Both paths start from the same structured source brief rather than a blank page, so the manual time reflects assembling and writing from provided material rather than composing from nothing. The AI portion replaces the bulk of that drafting effort, while the add-details step is identical on both sides, so the net AI savings come entirely from the drafting portion.

The quality of AI-drafted documents in our testing is suitable for downstream polish, capturing the core structure, key facts, and a serviceable first-pass voice. The result is not final without human refinement, but it consistently saves the writer the time and cognitive energy required to produce a clean first draft. For a mobile professional turning around a deliverable between meetings, that saved drafting time is often the difference between sending it today and sending it tomorrow.

AI-assisted drafting **cuts document production time by 42%** on a modern AMD laptop, freeing the writer from the cognitive load of the blank page while preserving the human polish step that matters most.



Workflow Total Time Savings (minutes)





Create Meeting Notes

Meeting summarization is one of the most time-consuming recurring tasks in commercial knowledge work, performed multiple times per week by most mobile professionals. It is also a task category that current AI models are particularly well-tuned for. The work involves taking a meeting transcript and producing a structured, readable summary with decisions, action items, and follow-ups. Manual production typically requires reviewing the transcript or recording, taking notes, and reformatting them into a shareable document, often squeezed in between the meeting that just ended and the next one starting.

Our test uses a recorded meeting transcript and asks an LFM2 2.6B local model to produce a structured summary with decisions, action items, and key discussion points. The manual workflow simulates a knowledge worker reviewing the same transcript and producing the same summary by hand. Both approaches target an identical deliverable, a clean meeting notes document suitable for distribution to stakeholders who did not attend.

AI-Enabled Steps	Manual Steps
The AI version uses a small local model specifically tuned for transcript summarization to produce structured meeting notes.	The manual version reads the same transcript and produces the same notes by hand, including formatting and cleanup.
Load meeting transcript	Load meeting transcript
Launch local LLM (LFM2 2.6B) via LM Studio	Read transcript carefully
Submit transcript with structured summary prompt	Take notes by hand
Receive AI-generated notes	Reformat notes into structured document
Light review and formatting	Final cleanup and formatting

REAL PRODUCTIVITY AND ALL-DAY BATTERY
FOR THE ROAD WARRIOR

Create Meeting Notes

The results are the most eye-opening in this report. The AI-enabled workflow completes in 1.6 minutes, compared to 28.5 minutes manually, a 94% reduction in time and a 17.8x speed-up. This gain sits well outside the 1.7x to 3.4x range seen on the other four workflows. The reason is that summarization plays directly to two strengths of modern small models, structure extraction and rewording from long input passages. The same characteristics that limit small models on open-ended reasoning tasks (limited reasoning depth, occasional factual drift) matter much less for summarizing structured conversational text where the source material is already present.

For a mobile professional attending three meetings per day, the cumulative time saved on note creation alone is more than six hours per week. Meeting summarization is by itself a credible justification for the AI-capable laptop refresh, even before the other workflows are considered. Beyond raw time savings, AI-generated notes are more consistent in structure and less prone to gaps caused by divided attention during the meeting itself, since the summarization step happens against a complete transcript rather than partial real-time notetaking on a laptop between other tasks.

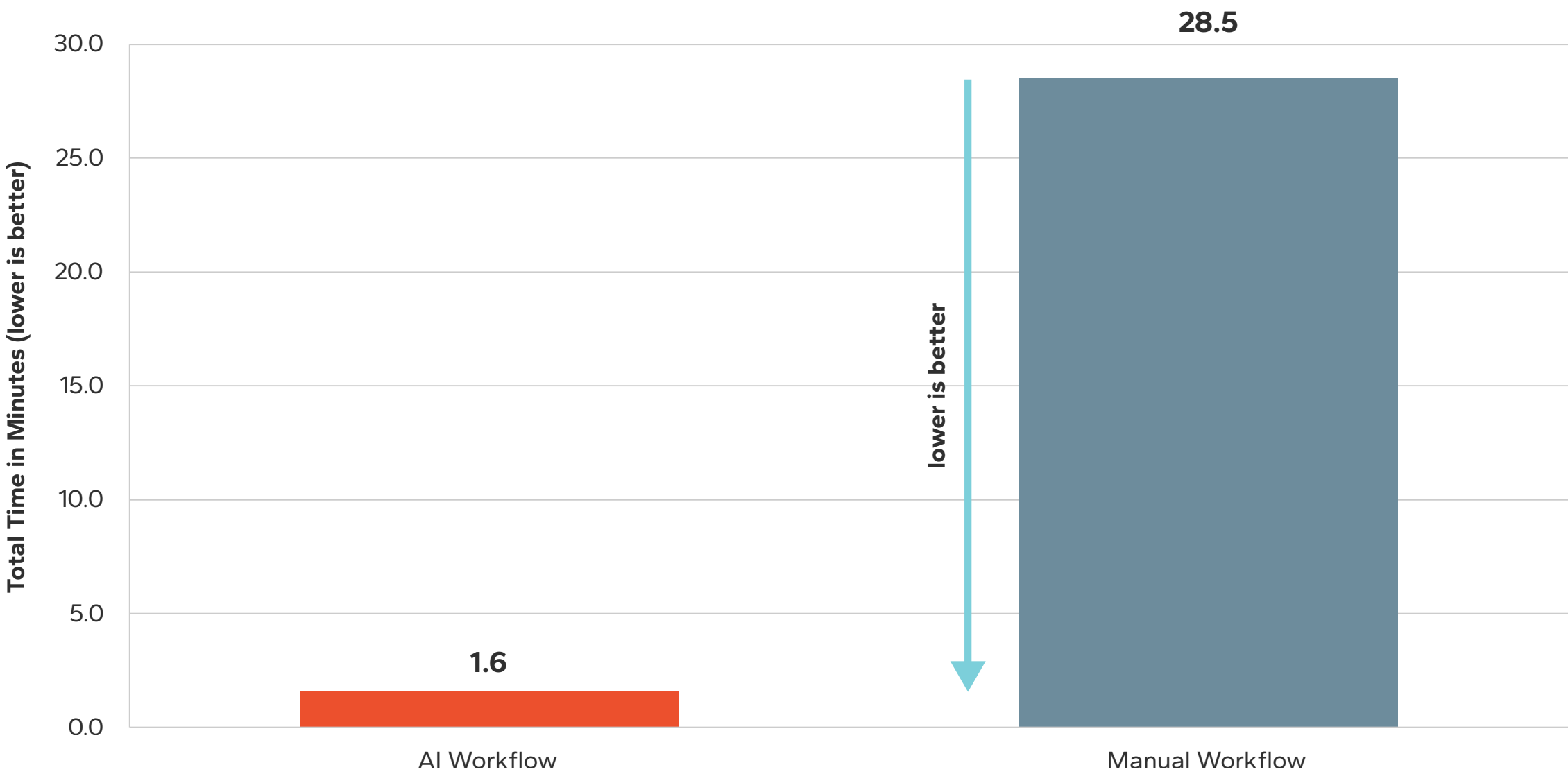
Workflow Total Time Savings (minutes)



Meeting summarization **runs 17.8x faster** with a local AI model on a modern AMD laptop, recovering more than six hours per week for the typical mobile professional.

AI vs. Manual Workflow - Road Warrior - Meeting Notes

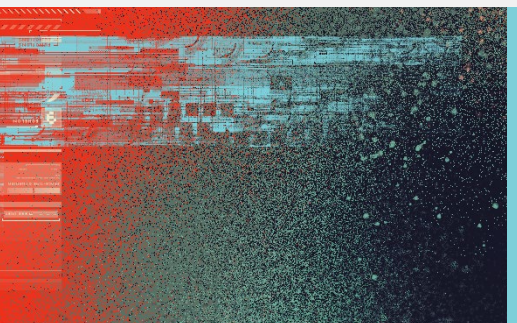
HP EliteBook X G2a powered by AMD Ryzen AI 7 PRO 450





Click-to-Do Slide Update

Updating a slide with refreshed content represents a common multi-application workflow that requires coordinating information across PowerPoint, a source document, and image editing. Our Click-to-Do scenario tests Windows Click-to-Do, a Copilot+ PC feature that allows the user to point at on-screen content and invoke AI actions directly without switching applications, a particular advantage on a laptop where screen space is limited.



The AI workflow uses Click-to-Do to select an image in PowerPoint and apply background blur, then selects source text in a Word document and instructs Click-to-Do to summarize it into bullet points, which are then pasted into the slide. The manual workflow performs the same operations the traditional way, switching between PowerPoint, Paint, and Word, and copying and editing each element by hand.

AI-Enabled Steps	Manual Steps
The AI version uses Click-to-Do, a Windows Copilot+ feature, to act on on-screen content across applications without manual context switching.	The manual version performs the same operations by switching applications and editing each element by hand.
Open PowerPoint deck and source Word document	Open PowerPoint deck and source Word document
Use Click-to-Do to select image in slide, apply background blur	Copy image from PowerPoint, switch to Paint to edit, paste back
Use Click-to-Do to select text in Word doc, summarize into bullets	Switch to Word, read and mentally summarize content
Paste AI-generated bullets into PowerPoint	Switch back to PowerPoint, type summary into text box
Light formatting cleanup	Adjust formatting manually

REAL PRODUCTIVITY AND ALL-DAY BATTERY
FOR THE ROAD WARRIOR

 Click-to-Do Slide Update

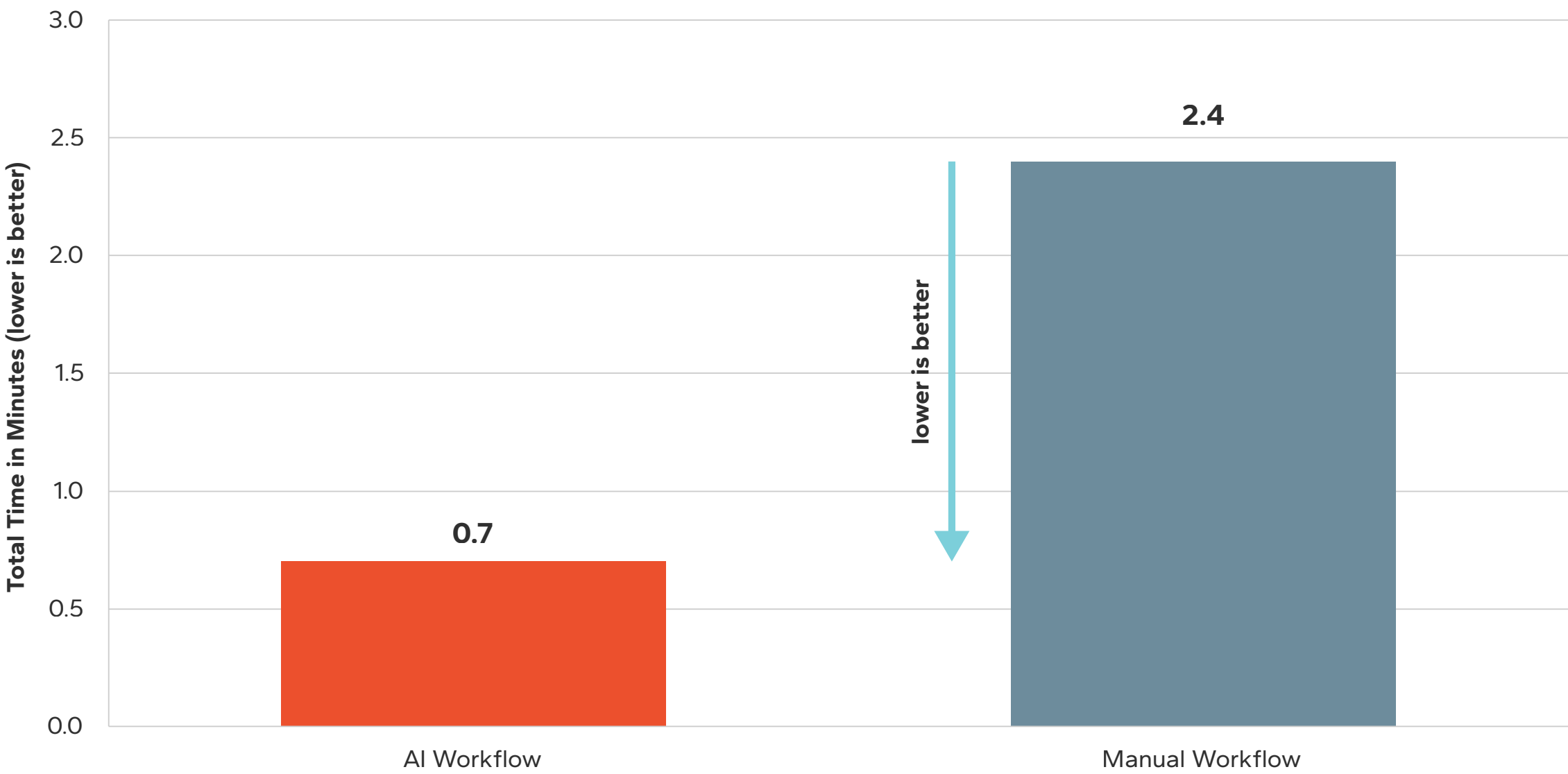
The AI workflow completes in 0.7 minutes versus 2.4 minutes for the manual approach, a 71% reduction in time and a 3.4x speed-up. In addition to the raw time savings, the AI workflow reduces context switches from seven application transitions to two, which reduces cognitive disruption. Multi-application workflows benefit twice from AI assistance, once from automation of the work itself, and again from staying in a single visual context, which is especially valuable on a smaller laptop display.

It is worth noting that Click-to-Do is gated by the Copilot+ PC requirements, which include integrated NPU performance of at least 40 TOPS. The prior-generation Ryzen 5 PRO 6650U system tested in this study does not meet those requirements, so this workflow was tested only on the new Ryzen AI 7 PRO 450 laptop. Some AI experiences are not available at all without modern silicon, and are not addressable through any amount of cloud-side processing because the experience itself depends on the on-device feature being present.

Click-to-Do **completes multi-application slide updates 3.4x faster** on a Copilot+ AMD laptop while cutting context switches from seven to two, a workflow the prior-generation laptop cannot run at all.

AI vs. Manual Workflow - Road Warrior - Click-to-Do Slide Update

HP EliteBook X G2a powered by AMD Ryzen AI 7 PRO 450



Workflow Total Time Savings (minutes)





Research Client Documents

Reviewing and extracting insight from client documents is a core task for the consultant, the sales engineer, and any mobile professional who walks into a meeting needing to know what a long document says. The work involves reading through source material, identifying the relevant points, and producing answers or a synthesis. Our Research Client Documents scenario asks a Qwen3.5 4B local model to analyze a set of client documents and answer questions against them, compared to a professional doing the same reading and synthesis by hand.

This workflow is a strong fit for the road warrior profile specifically because it so often happens under time pressure, in the minutes before a client meeting or between sessions at a conference. Offloading the first-pass reading and extraction to a local model lets the professional spend their limited time verifying and applying the insight rather than hunting for it.



AI-Enabled Steps	Manual Steps
The AI version uses a small local model to analyze client documents and answer questions against the source material.	The manual version reads each document and produces the same answers by hand.
Load client documents into the workflow	Load client documents into the workflow
Launch local LLM (Qwen3.5 4B) via LM Studio	Read each document carefully
Submit documents with analysis and question prompt	Identify relevant points and passages
Receive AI-generated answers and synthesis	Compile answers and synthesis
Verify answers against source	Review for completeness

REAL PRODUCTIVITY AND ALL-DAY BATTERY
FOR THE ROAD WARRIOR

 Research Client Documents

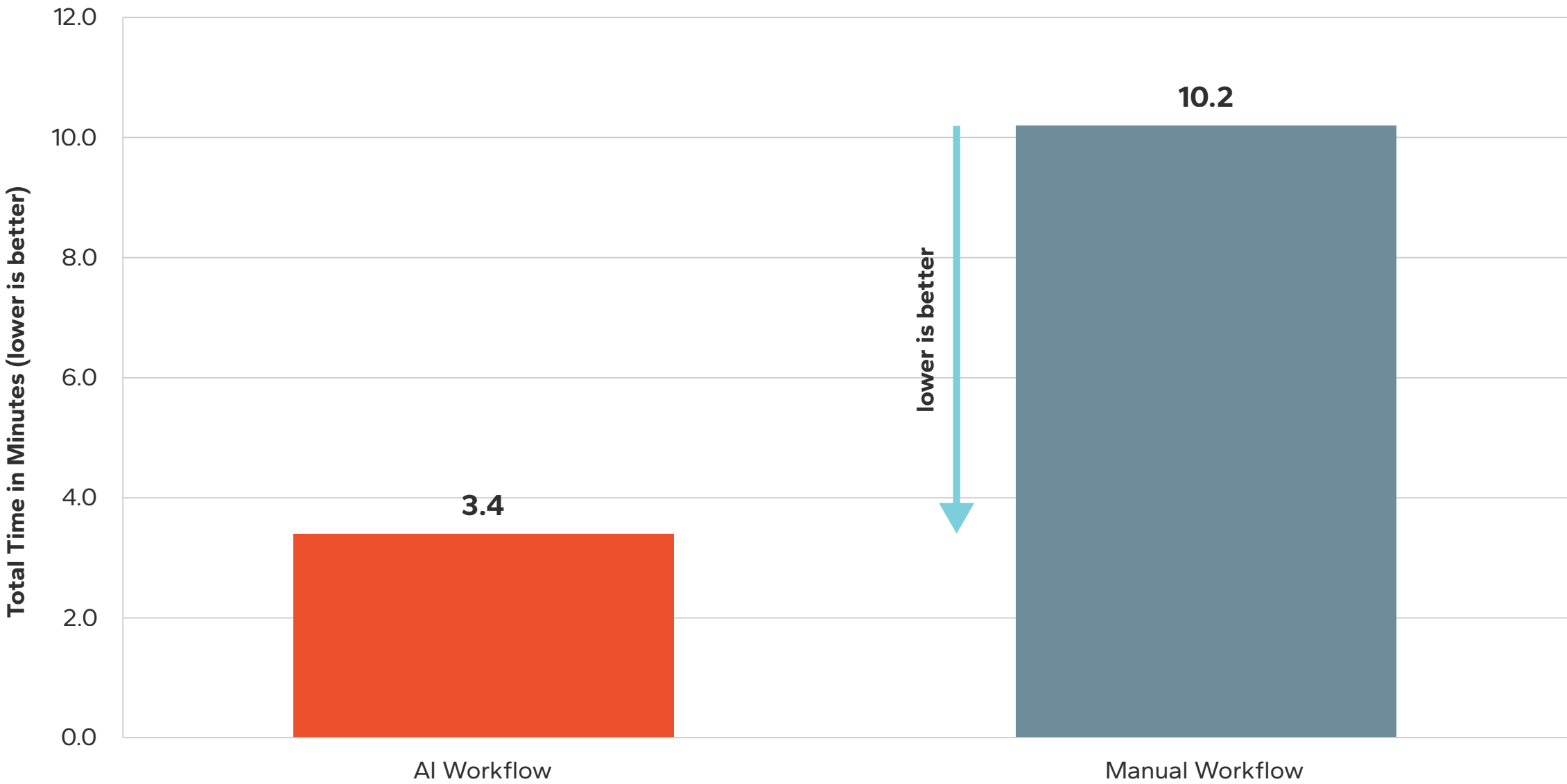
On the new Ryzen AI 7 PRO 450 laptop, the AI workflow completes in 3.4 minutes versus 10.2 minutes manually, a 67% reduction in time and a 3.0x speed-up. The same workflow on the prior-generation Ryzen 5 PRO 6650U takes 5.2 minutes, still an AI win but a smaller one, and 1.5x slower than the new machine. The new silicon does not change whether AI helps here, it changes how much, turning a useful tool into a genuinely fast one.

Document analysis is a workload where the quality of the on-device model matters as much as the speed. In our testing the local model reliably surfaced the relevant passages and produced a serviceable synthesis, which the professional then verified. As with the other workflows, the AI output is a starting point rather than a final answer, but it consistently compresses the slowest part of the task, the initial read, into a fraction of the time.

AI-assisted document research **runs 3.0x faster** on a modern AMD laptop, compressing the slow initial read of client material into minutes, exactly when the road warrior has the least time to spare.

AI vs. Manual Workflow - Road Warrior - Research Client Documents

HP EliteBook X G2a powered by AMD Ryzen AI 7 PRO 450



Workflow Total Time Savings (minutes)

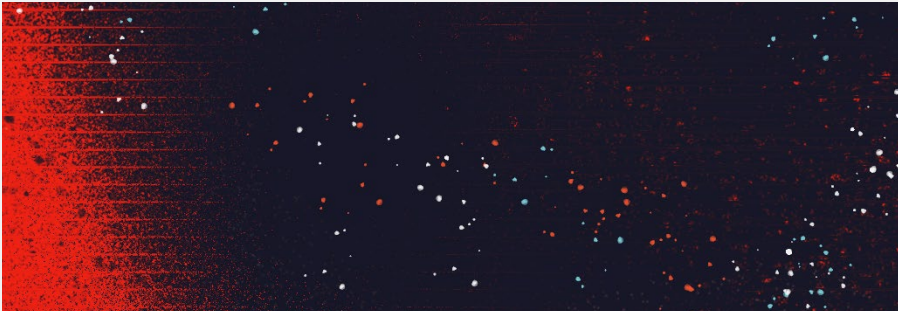




Generate Tasks from Emails

Knowledge workers regularly need to extract actionable items from a thread or batch of emails. The work involves reading through messages, identifying what requires action, and producing a structured list of tasks with owners and deadlines where possible. Our test asks a Qwen3.5 4B local model to perform that extraction from a sample email batch, compared to a knowledge worker doing the same work by hand.

AI-Enabled Steps	Manual Steps
The AI version uses a small local model to extract structured task items from a batch of emails.	The manual version reads each email and extracts task items by hand.
Load email batch into the workflow	Load email batch into the workflow
Launch local LLM (Qwen3.5 4B) via LM Studio	Read each message carefully
Submit emails with task-extraction prompt	Identify task items as they appear
Receive structured task list	Compile structured task list
Light verification of extracted tasks	Review the list for completeness



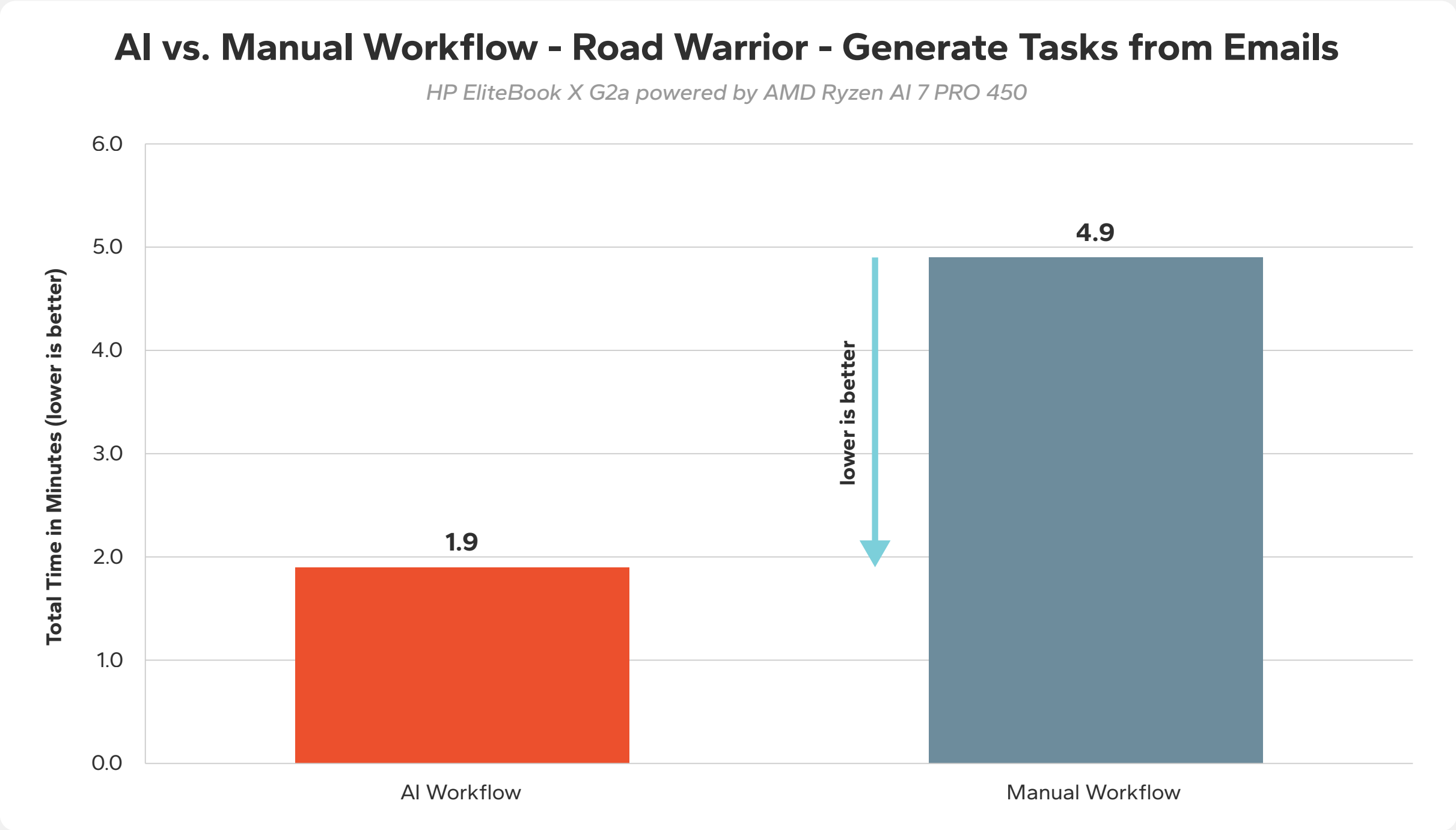
REAL PRODUCTIVITY AND ALL-DAY BATTERY
FOR THE ROAD WARRIOR

 Generate Tasks from Emails

On the new Ryzen AI 7 PRO 450 laptop, the AI workflow completes in 1.9 minutes versus 4.9 minutes manually, a 61% reduction in time and a 2.6x speed-up. The same workflow on the prior-generation Ryzen 5 PRO 6650U takes 3.6 minutes, a 1.9x slower result than the new machine but still faster than doing the work by hand.

For the road warrior triaging a full inbox after a flight, the value is both the time saved and the reduction in dropped commitments. A structured task list generated in under two minutes surfaces action items that are easy to miss when skimming a crowded inbox on a small screen between other tasks.

AI task extraction **runs 2.6x faster** on a modern AMD laptop and remains a clear win even on the prior generation, turning a full inbox into a structured task list in under two minutes.



Workflow Total Time Savings (minutes)



Implications for Road Warriors with AI Laptops

Our testing shows that AI-capable AMD laptops can deliver clear and measurable productivity advantages for mobile professionals. Across real-world workflows spanning content drafting, meeting summarization, multi-application UI interactions, document research, and structured extraction, AI-enabled workflows on the new AMD Ryzen AI 7 PRO 450 laptop completed in 17.4 minutes a representative bundle of tasks that took 62.8 minutes manually, a 72% reduction. For mobile professionals whose daily output is often constrained by the speed and battery of their primary work surface, that difference represents a meaningful expansion of productive capacity, and it holds up away from a charger.

The most important finding for a mobile professional is where the AI runs. On this laptop, moving the local model from the GPU to the integrated NPU raised the office-productivity score while halving the battery drawn, and the new laptop as a whole delivers 1.8x the productivity of the prior generation on just over a third as much battery capacity under multitasking. This is also the first of our studies to measure not just how fast AI runs on new versus old silicon but what it costs in battery. New silicon is not just faster, at 1.9x on the same AI workload, it does more work per watt. For a workforce that increasingly works on the move, that efficiency is not a footnote, it is the headline.

The productivity and battery gains shown here are available now, but only on hardware specified for them. Refresh cycles that do not prioritize AI silicon and on-battery efficiency will leave the largest single productivity opportunity of the current mobile generation on the table. From an IT decision-maker perspective, an NPU-equipped fleet also carries option value beyond the workflows measured in this study. Endpoint security and device management vendors are beginning to build NPU-accelerated capabilities, and the on-device AI feature set in Windows continues to expand with each release, so a laptop specified with an NPU today is positioned for capabilities that have not shipped yet. Those capabilities are outside the scope of our testing, but they are a reasonable input to a refresh decision.

Implications for Road Warriors with AI Laptops

At the same time, AI adoption introduces operational considerations that should be addressed during planning rather than after wide system deployment. Local-model inference, where this report has shown the largest gains, keeps data on the device and avoids the cloud egress and licensing costs that come with cloud-only AI workflows. That is a security and cost benefit, and on a mobile device it is also a resilience benefit, since the work does not depend on connectivity. It does, however, require hardware capable enough to deliver acceptable latency and battery life. Policies around model selection, model updates, prompt logging, and acceptable-use guidance should be developed in parallel with any hardware rollout like this. The investment is in the combined hardware-software-policy package, not just the laptops themselves.

Ultimately, the case for AI-capable AMD laptops for road warriors is a practical one. The productivity and battery gains shown here can justify the investment, and the rate of improvement across models and software suggests that future benefits will be even larger. Enterprises that enable AI-enhanced mobile work now are not just optimizing the most time-constrained and power-constrained segment of their workforce, they are building the foundation for a workforce that can act faster from anywhere, communicate more clearly, and adapt more effectively to the increasing pace of business in the coming years.

Key Highlights:



Running AI on the integrated NPU delivers **11% higher productivity on 50% less battery.**



Workflow steps like meeting summarization run up to **17.8x as fast** with AI than manual work.



The new AMD Ryzen AI laptop **reduces task time by 72%** across our road warrior workflows.

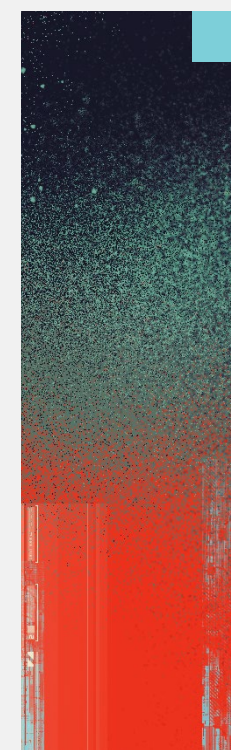


The new Ryzen AI CPU runs the **same AI workload 1.9x as fast** as the prior-generation AMD laptop while **drawing just a third of the battery** capacity.

Appendix: Our Testing Theory

Developing a fair and comprehensive comparison between AI-enabled and traditional workflows requires careful consideration of methodology, metrics, and real-world applicability. Our approach centers on creating reproducible, realistic workflows that mirror actual user behavior across different personas. Rather than constructing artificial benchmarks that might favor one approach over another, we have developed complete task sequences that encompass the full complexity of modern work.

For each scenario we test, we create two distinct implementations: one that leverages currently available AI capabilities and another that follows traditional processes. Both workflows target identical outcomes and deliverables, ensuring that we are measuring different paths to the same destination rather than comparing fundamentally different tasks. This approach acknowledges that AI-enabled workflows may use entirely different applications or involve novel steps that have no traditional equivalent, such as automated content generation followed by human refinement.



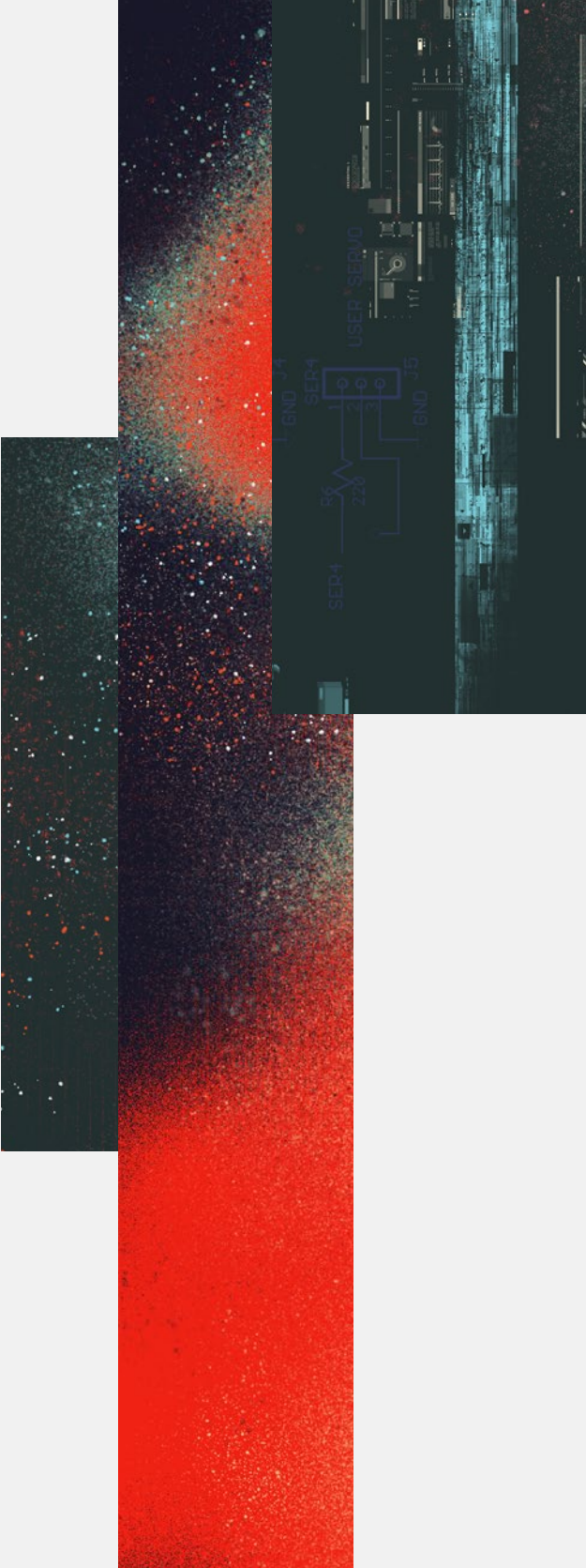
In this study we extended our methodology to include two additional variables that matter specifically for the mobile professional. The first is a second hardware configuration. Each AI-enabled workflow was run on two AMD laptop systems, a current-generation Ryzen AI 7 PRO 450 and a prior-generation Ryzen 5 PRO 6650U, using the same AI models, prompts, and input materials on both, so that the silicon contribution can be isolated independently of the model and workflow design. The second is battery and multitasking. We ran an office-productivity benchmark concurrently with a local AI workload and recorded both the productivity score and the battery capacity drawn in milliwatt-hours, separately for the NPU and GPU execution paths and for the new machine against the old. This is what allows the study to report not just the speed of on-device AI but its energy cost.

Our testing philosophy embraces the current reality of hybrid AI deployment. Despite the rapid evolution of AI PCs and the trajectory toward local, on-device processing, today's optimal solutions often combine cloud-based and local compute resources. Four of the five workflows in this study use local large language models running entirely on the laptop. The fifth uses Windows Copilot+ Click-to-Do, which is a Windows feature with its own internal hybrid execution model. We have documented where local versus cloud processing applies for each workflow, allowing readers to make informed decisions about security, privacy, and infrastructure requirements.

Appendix: Our Testing Theory

The metrics we capture extend well beyond time measurements, though efficiency does form the quantitative basis of our analysis. For local-LLM workflows we also captured Time-to-First-Token and generation throughput, which together explain the source of the silicon contribution measured in the workflow timings, and under multitasking we captured battery capacity drawn. These metrics are reported alongside the workflow times in the Hardware and Efficiency Story section.

Our testing infrastructure combines automated scripts with manual execution to ensure both reproducibility and realism. We use AutoHotKey scripts to standardize user interactions, eliminating variability from factors like typing speed or navigation patterns. Python automation handles data collection, timing measurements, and result aggregation. Specialized tools, including Camo Studio for screen recording and playback, LM Studio for local large language model deployment, and VB Cable for audio routing, ensure comprehensive data capture. Each workflow undergoes a minimum of three iterations, with results averaged to account for system variability and ensure statistical validity. All reported figures are rounded for presentation, with derived values computed from those rounded figures so that the tables and totals reconcile.



Important Information About this Report

Contact Information

Signal65 | signal65.com | info@signal65.com

Contributors

Ryan Shrout

President & GM - Signal65

Ken Addison

Client Performance Director - Signal65

Ken Fetcher

Performance Analyst and Software
Engineer - Signal65

Inquiries

Contact us if you would like to discuss
this report and Signal65 will respond
promptly.

Citations

This paper can be cited by accredited
press and analysts, but must be cited
in-context, displaying author's name,
author's title, and "Signal65." Non-press
and non-analysts must receive prior
written permission by Signal65 for any
citations.

Licensing

This document, including any supporting
materials, is owned by Signal65. This
publication may not be reproduced,
distributed, or shared in any form without
the prior written permission of Signal65.

Disclosures

Signal65 provides research, analysis,
advising, and lab services to many high-tech
companies, including those mentioned in
this paper. Research of this document was
commissioned by AMD.

In Partnership with:

AMD 
together we advance_

About Signal65

Signal65 exists to be a source of data in
a world where technology markets and
product landscapes create complex and
distorted views of product truth. We strive
to provide honest and comprehensive
feedback and analysis for our clients in order
for them to better understand their own
competitive positioning and create optimal
opportunities to market and message their
devices and services.



REAL PRODUCTIVITY AND ALL-DAY BATTERY
FOR THE ROAD WARRIOR

System Configurations

	HP ELITEBOOK X G2A	HP ELITEBOOK 845 G9
CPU	AMD Ryzen AI 7 PRO 450	AMD Ryzen 5 PRO 6650U
Graphics	AMD Radeon™ 860M	AMD Radeon™ 660M
RAM	32GB LPDDR5X-8533	32GB DDR4-4800
Storage	1TB Samsung PM9C1	1TB Western Digital SN730
NPU	AMD XDNA 2	None
Display	14” 1920x1200	14” 1920x1200
System BIOS	Y88 01.01.18	U82 01.15.00
Operating System	Windows 11 Pro 26200.8655	Windows 11 Pro 26200.8655
Virtualization Based Security	Enabled	Enabled

Applications Used

Microsoft Office 2604	Lemonade Server 10.3.0
Microsoft Teams 2612.3106.4722.3411	LMStudio 0.4.16
UL Procyon 2.11.2513	
AnythingLLM v1.14.0	



signal**65**