

M C P   S E R V E R

N O   C O D E

C L O U D   H O S T E D

# assignment-time-per-page is an academic pacing MCP.

Measure exactly how fast you move through your research and writing.

assignment-time-per-page gives your AI client the ability to track and analyze your academic workflow. It calculates how much time you spend on every page, compares different study sessions, and checks if your current pace matches your goals for drafting or research tasks.

**A+** Quality Score 100/100

academic

pacing

efficiency

workflow

measurement



# The connectivity layer between AI and the world's software.

---

Vinkius sits between AI and every application. All communication passes through Vinkius Cloud via the Model Context Protocol (MCP) — with governance, observability, and security at every layer.

# Your AI Connections Run Through Vinkius Cloud

The world's largest  
managed MCP catalog

Vinkius is the connectivity layer where AI connects to the software your business already runs. We handle the hosting, the security, the credentials, the uptime — you get agents that actually do things.

We operate the world's largest managed MCP catalog. Major SaaS platforms, CRMs, databases, and cloud providers — running, monitored, production-ready. This MCP server is hosted and maintained by the Vinkius Cloud for AI Agents.

*The agent doesn't manage credentials, doesn't manage uptime, doesn't manage security. Vinkius does.*

— Architecture principle

---

## Four Pillars of the Vinkius Runtime

### 01 — Security by design

Credentials stay encrypted at rest via AES-256. The AI agent never touches raw keys — they're injected into a sandboxed V8 isolate at runtime. Actions are logged, and connections have an emergency kill switch.

### 03 — Deterministic observability

Eight immutable metrics per endpoint: request volume, p95 latency, error rate, active connections, cost attribution. A live payload feed logs every tool call with mutation detection.

### 02 — Built on MCP Fusion

This MCP server was built with **MCP Fusion**, the open-source framework (Apache 2.0) that powers the entire Vinkius catalog. Schema-as-firewall strips undeclared fields, compiled PII redaction runs at zero overhead, and cryptographic lockfiles produce git-diffable audit trails.

### 04 — Autonomous operations

Servers are deployed, monitored, and patched autonomously. New capabilities and security patches ship weekly. Zero-downtime deployments ensure continuous availability across all managed MCP servers.

**AES-256**

Encryption at rest

**Ed25519**

PKI vault signatures

**24h TTL**

Ephemeral session keys

**V8 Isolate**

Sandboxed execution

---

## One Token. Instant Access.

Every MCP server on Vinkius is accessed through a **Connection Token**. Tokens are generated in the cloud dashboard and produce a unique MCP endpoint URL. Paste this URL into any MCP-compatible client — no SDK required.

A single token can serve **multiple AI clients simultaneously**, or you can issue separate tokens per client for granular access control. Each token tracks its own request count, last activity timestamp, and can be individually enabled or revoked.

**MCP ENDPOINT**`https://edge.vinkius.com/{token}/mcp`

Claude



Cursor



VS Code



Windsurf



Grok



Gemini

---

## Security Is the Architecture

Security in Vinkius is not a feature — it's the foundation of the runtime. The gateway enforces multiple independent protection layers between AI agents and third-party APIs.

**01 — Ed25519 PKI Vault**

Every workspace has an Ed25519 Master Key. Session keys are generated ephemerally (24h TTL) and signed by the Master Key. Credentials never leave the vault boundary.

**02 — V8 Isolate Sandboxing**

Tool code runs inside isolated-vm V8 isolates with 64 MB memory caps and per-request timeouts. No filesystem access, no network access except through the SSRF-guarded fetch bridge.

**03 — SSRF Guard**

All outbound HTTP requests are DNS-resolved and validated before execution. Private IP ranges (10.x, 172.16-31.x, 192.168.x, AWS metadata 169.254.x) are blocked at the network layer.

**05 — Cryptographic Audit Trail**

Every request is signed into a SHA-256 hash chain with Ed25519 signatures. Events form a tamper-proof, SIEM-exportable forensic record.

**04 — DLP & PII Redaction**

A ResponseGuard pipeline intercepts every tool response. Configurable redaction patterns strip sensitive fields (emails, SSNs, card numbers) before data reaches the AI agent.

**06 — Honeypot Trap System**

Phantom credentials are injected into isolated environments. If a honeypot is used outside Vinkius infrastructure, the server is quarantined instantly.

## Emergency Kill Switch

EU AI Act Art. 14(1)  
Compliant

The kill switch is an **emergency halt** mechanism — not a simple toggle. When triggered, it executes three actions atomically:

**01 — Server deactivated**

The MCP server is immediately taken offline across the entire cluster.

**02 — All tokens revoked**

Every connection token is invalidated. Total lockout — reconnection blocked until new tokens are issued.

**03 — WebSocket connections killed**

Active connections terminated via Redis pubsub broadcast. Propagates to every runtime node in the cluster.

## Full Visibility. Zero Guesswork.

The Vinkius cloud dashboard includes a full MCP Governance suite — real-time analytics and security controls for production AI operations.

**Control Plane**

KPI dashboard with request volume, latency, success rate, token consumption, and AI-generated operational briefings.

**FinOps**

Cost tracking per tool, payload compression savings, budget optimization signals, and consumption trends.

**Firewall & DLP**

PII redaction activity, sensitive data protection counters, and security event timeline.

**Agent Activity**

Which AI clients are connecting, how often, and what they're doing — real-time session tracking.

**Tool Health**

Slowest and most error-prone tools, with actionable root-cause insights and performance baselines.

**Incident Log**

Error trends, failure rates, status-code breakdowns, and forensic audit trail access.

Get started at [cloud.vinkius.com](https://cloud.vinkius.com) — connect your AI agent in under 60 seconds.

# assignment-time-per-page MCP

4 tools available

Cloud-hosted on Vinkius

Stop guessing how long your writing sessions actually take. This MCP gives your agent the math it needs to analyze your academic productivity. Instead of vague feelings about being slow or fast, you get hard numbers on your minutes per page. You can use it to see if your research phase is dragging or if your drafting speed is hitting the right marks. It works by breaking down your total time against your page count to give you a clear picture of your efficiency. You can compare how you worked on a paper last week versus how you are working today to see if your workflow is actually improving. It is a practical way to keep your academic output consistent and predictable.

---

## Core Capabilities

### 01 — Page-level math

Your agent calculates the specific time cost for every page you complete.

### 02 — Session comparison

Your agent compares the speed of different work sessions to find patterns.

### 03 — Pacing validation

Your agent checks if your current speed matches your predefined productivity limits.

### 04 — Workload summaries

Your agent generates qualitative reports on how you are handling your workload.

# One Click on Vinkius — From Prompt to Execution

Available at [vinkius.com/en/ai-agent-connect/assignment-time-per-page](https://vinkius.com/en/ai-agent-connect/assignment-time-per-page) — connect your AI agent in three steps.

- 01 Connect the MCP to Claude, Cursor, or Windsurf via Vinkius.
- 02 Provide your agent with your total time spent and pages completed.
- 03 Ask the agent to calculate your efficiency or compare sessions.
- 04 Receive specific pacing data and summaries directly in your chat.

Connect the MCP to your client and start feeding it your session data.

## Built For

This MCP is built for anyone managing heavy academic or research workloads who needs to track output speed.

**Researchers**

Hand off your session durations and page counts to get efficiency metrics.

**Students**

Give your agent your writing times to monitor your study pacing.

**Academic Writers**

Use the tool to compare drafting speeds across different manuscripts.

## What Changes When You Connect

- 01 Converts total time and page counts into minutes per page.
- 02 Identifies if your current pace is too slow or too fast for your goals.
- 03 Provides a way to compare efficiency between different work sessions.

## 04 Generates qualitative summaries of your academic workload.

# Real-World Applications

### Drafting Speed Checks

Check if your initial drafting phase is moving at a sustainable pace.

### Research Efficiency

Compare how much time you spend reading and processing pages during research.

### Workflow Comparison

See if a new writing environment or method actually makes you faster.

### Productivity Monitoring

Validate that your current assignment pace stays within acceptable bounds.

# Patterns to Avoid

## Vague productivity guesses

### ❌ AVOID

You tell your agent that you feel like you are writing slowly today without providing any specific time or page data.

### ✓ INSTEAD

Provide your actual session durations and page counts to ``get_page_efficiency`` for a precise calculation of your minutes per page.

## Ignoring pace shifts

### ❌ AVOID

You continue a research phase that is taking twice as long as usual without checking if you are still on track.

### ✓ INSTEAD

Use ``check_productivity_thresholds`` to see if your current speed is drifting outside your expected productivity bounds.

## Comparing unrelated work

### ❌ AVOID

You try to judge your research speed by comparing it directly to your final drafting speed.

### ✓ INSTEAD

Use ``compare_assignments`` to look at efficiency levels between two similar types of assignment sessions instead.

## The Right Fit

Connect this MCP if you manage heavy research or academic workloads and need hard numbers to track your output. Do not connect it if you just want a general timer or if your work doesn't involve measurable page counts. The decision depends on whether you need to validate your drafting speed against specific productivity limits.



# assignment-time-per-page 4 Tools

These tools turn your raw session data into actionable metrics for tracking academic output and drafting speed.

| #  | TOOL                          | DESCRIPTION                                                                               |
|----|-------------------------------|-------------------------------------------------------------------------------------------|
| 01 | check_productivity_thresholds | This tool checks if your current assignment pace stays within your expected bounds.       |
| 02 | compare_assignments           | Use this to compare the efficiency levels of two different assignment sessions.           |
| 03 | get_pacing_summary            | This tool provides a qualitative summary of your workload based on your efficiency data.  |
| 04 | get_page_efficiency           | This tool calculates the exact number of minutes you spent on each page of an assignment. |

---

## See It in Action

Real prompts you can use once this MCP is connected to your AI agent through Vinkius Cloud.

**U** I spent 120 minutes writing 4 pages. How much time did I spend per page?



You spent 30 minutes per page.

**U** I finished 10 pages in 150 minutes. What is my pacing category?



Your pacing is 15 minutes per page, which falls into the Standard category.

**U** Compare two sessions: Session A (60 mins, 2 pages) and Session B (100 mins, 5 pages).



Session B was more efficient, with a difference of 10 minutes per page.

---

## Frequently Asked Questions

### 01 What can this MCP do for my writing?

It calculates exactly how many minutes you spend per page and tells you if your pace is within your expected bounds.

### 02 Which AI clients can use this MCP?

You can use this with any MCP-compatible client like Claude, Cursor, or Windsurf.

### 03 Can I compare two different writing sessions?

Yes, you can use the comparison tool to see which session was more efficient.

**04 How does it define my productivity?**

It uses your time and page counts to generate a pacing summary and check against your thresholds.

---

**05 Do I need to host this myself?**

No, Vinkius hosts and manages the MCP for you so it is ready to use immediately.

---

**06 How do I calculate my writing speed?**

You can use the ``get_page_efficiency`` tool by providing the total minutes spent and the number of pages completed.

---

**07 Can I compare two different study sessions?**

Yes, the ``compare_assignments`` tool allows you to compare the efficiency of two different sessions to see which was faster.

---

**08 What modes are supported for checking productivity?**

The system supports 'drafting', 'editing', and 'researching' modes via the ``check_productivity_thresholds`` tool.

---

# Go Live in 60 Seconds

Get your connection token from [cloud.vinkius.com](https://cloud.vinkius.com), then paste the endpoint URL into any MCP-compatible client.

YOUR MCP ENDPOINT

`https://edge.vinkius.com/[TOKEN]/mcp`

| CLIENT                                                                                              | WHERE TO CONFIGURE                                                                                         |
|-----------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------------------------------------|
|  <b>Claude AI</b>  | Profile → Customize → Connectors → "+" → Add custom connector → Paste endpoint                             |
|  <b>Cursor</b>     | Settings → Features → MCP Servers → "+ Add New MCP Server" → Type: SSE → Paste endpoint                    |
|  <b>VS Code</b>  | Ctrl/Cmd+Shift+P → "MCP: Add Server" → add <code>"assignment-time-per-page": {<br/>  "url": "..." }</code> |
|  <b>Windsurf</b> | MCP Settings → <code>mcp_settings.json</code> → Add endpoint URL                                           |
|  <b>ChatGPT</b>  | Settings → Tools & plugins → Add MCP server → Paste endpoint                                               |
|  <b>Gemini</b>   | Extensions → Add MCP Server → Paste endpoint URL                                                           |

ASK AN AI ABOUT THIS

Let your preferred AI explain this MCP server

-  Ask ChatGPT 
-  Ask Claude 
-  Ask Perplexity 
-  Ask Gemini 
-  Ask Grok 

READY TO CONNECT

# assignment-time-per-page is live on Vinkius Cloud.

---

Get your connection token, paste it into your AI agent, and  
start building. No SDK. No deployment. Just results.

**Start at [cloud.vinkius.com](https://cloud.vinkius.com) →**

[vinkius.com](https://vinkius.com) · [support@vinkius.com](mailto:support@vinkius.com)

INDEPENDENT PLATFORM DISCLAIMER

Vinkius is an independent platform and is not affiliated with, endorsed by, sponsored by, verified by, or otherwise authorized by assignment-time-per-page. All third-party trademarks, logos, and brand names are the property of their respective owners. Their use in this document is strictly for informational purposes to identify service compatibility and interoperability.

DOCUMENT INFORMATION

|            |                                      |
|------------|--------------------------------------|
| Generated  | September 2026                       |
| MCP Server | assignment-time-per-page MCP         |
| Server ID  | 01a0cc7a-4379-72d5-81e7-a8ffe8d2e846 |
| Platform   | Vinkius Cloud for AI Agents          |
| Endpoint   | https://edge.vinkius.com/{token}/mcp |

LICENSE & USAGE

This document is generated automatically by the Vinkius PDF Engine. Content reflects the MCP server configuration at the time of generation and may change as updates are deployed. For the most current information, visit [vinkius.com/en/ai-agent-connect/assignment-time-per-page](https://vinkius.com/en/ai-agent-connect/assignment-time-per-page).