

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

# Productivity Output Rate MCP

Turn your historical work data into actionable capacity plans.

Productivity Output Rate MCP lets your AI client calculate exactly how much work you get done per hour. Instead of guessing your capacity, you can use this tool to analyze performance trends, compare your results against regional benchmarks, and predict how much time future projects will actually take based on your real history.

**A+** Quality Score 100/100

efficiency

metrics

workload

performance

planning



# 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.

# Productivity Output Rate MCP

4 tools available

Cloud-hosted on Vinkius

Stop guessing how much work you can handle in a week. This MCP gives your AI agent the math it needs to understand your actual output. You can use it to find your baseline speed, see if your performance is improving or slipping over time, and compare your pace against standard benchmarks. It's built for anyone who needs to turn raw task data into a predictable schedule. When you have a new project on your plate, your agent can look at your past sessions to tell you exactly how many hours you'll need to finish. It moves you from vague estimates to data-backed planning.

---

## Core Capabilities

### 01 — Output Calculation

Your agent calculates items completed per hour for any given timeframe.

### 02 — Trend Analysis

The AI evaluates performance shifts across multiple recorded work sessions.

### 03 — Capacity Forecasting

Your agent predicts the time required for future volumes of work.

### 04 — Benchmark Comparison

The tool compares your personal efficiency against regional standards.

# One Click on Vinkius — From Prompt to Execution

Available at [vinkius.com/en/ai-agent-connect/productivity-output-rate](https://vinkius.com/en/ai-agent-connect/productivity-output-rate) — connect your AI agent in three steps.

- 01
- Connect your preferred MCP-compatible client to Vinkius.
- 02
- Provide your task and time data to your AI agent.
- 03
- Ask your agent to run specific tools like `get_output_rate`.
- 04
- Receive calculated rates, trends, or capacity estimates immediately.

Connect your client to Vinkius and start running math on your work data.

## Built For

This is for professionals who manage high volumes of repetitive tasks and need to predict their availability.

### Project Managers

They hand off task completion data to the AI to build more accurate project timelines.

### Operations Leads

They use the tool to compare team output against regional efficiency benchmarks.

### Freelancers

They use output rates to quote clients more accurately based on their actual speed.

## What Changes When You Connect

- 01
- Replaces guesswork with math-based workload planning.
- 02
- Identifies performance shifts through multi-session trend analysis.
- 03
- Provides regional benchmarks to contextualize your efficiency.

- 04 Predicts future time requirements using historical performance data.
- 

---

## Real-World Applications

### Project Scoping

Use your historical output rate to tell a client exactly how many days a new project will take.

### Resource Planning

Calculate the capacity requirement for a large batch of tasks to avoid overcommitting.

### Performance Tracking

Analyze your productivity trends to see if you are getting faster or slower over several weeks.

### Efficiency Audits

Compare your current output rate against efficiency benchmarks to see where you stand.

---

---

## Patterns to Avoid

### Guessing project timelines

#### ❌ AVOID

You tell a client a project will take three days based on a gut feeling. You end up overcommitting and missing the deadline.

#### ✓ INSTEAD

Use ``calculate_capacity_requirement`` to get a realistic estimate based on your actual historical speed.

---

### Ignoring performance shifts

#### ❌ AVOID

You assume you are working at the same speed you were six months ago. You fail to notice your output is actually dropping.

#### ✓ INSTEAD

Run ``analyze_productivity_trends`` to see how your performance changes across different work sessions.

---

### Using arbitrary efficiency goals

#### ❌ AVOID

You set a target speed that has no basis in reality or industry standards. You frustrate yourself by chasing impossible numbers.

#### ✓ INSTEAD

Check ``get_efficiency_benchmarks`` to see how your current pace compares to actual regional standards.

---

## The Right Fit

Connect this MCP if you manage high volumes of repetitive tasks and need to stop overcommitting. If you work on highly unpredictable, non-repetitive creative projects where every task is unique, the math won't be as reliable for you.

# Productivity Output Rate 4 Tools

These tools convert your historical task data into actionable math for scheduling and performance tracking.

| #  | TOOL                           | DESCRIPTION                                                                                                                               |
|----|--------------------------------|-------------------------------------------------------------------------------------------------------------------------------------------|
| 01 | analyze_productivity_trends    | This tool tracks how your output changes across different work sessions. It helps your agent spot patterns in your performance over time. |
| 02 | calculate_capacity_requirement | Use this to estimate the time needed for a specific workload. It uses your historical data to predict how long future tasks will take.    |
| 03 | get_efficiency_benchmarks      | This tool pulls standard values used to categorize efficiency levels. It lets you see how your pace compares to regional standards.       |
| 04 | get_output_rate                | This tool calculates your basic productivity rate for a set period. It tells you exactly how many items you complete per hour.            |

---

## See It in Action

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

**U** What is my output rate if I finished 20 tasks in 5 hours?



Your output rate is 4.0 items per hour, which is considered an Optimal efficiency level.

**U** How long will it take me to finish 50 tasks if my average rate is 5 items per hour?



It will take approximately 11 hours to complete 50 tasks, including a buffer for unexpected delays.

**U** Analyze my productivity trends for these sessions: [{"items": 10, "hours": 2}, {"items": 15, "hours": 2}, {"items": 12, "hours": 2}]



Your average rate is 6.17 items per hour, with a stable trend and low volatility.

---

## Frequently Asked Questions

### 01 How does this MCP calculate my productivity?

It uses the `get_output_rate` tool to divide the number of completed items by the time spent working.

### 02 Can I use this to plan future projects?

Yes, the `calculate_capacity_requirement` tool uses your past performance to estimate the time needed for future work volumes.

### 03 What clients can I use this with?

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

**04 Does it compare my speed to others?**

The `get_efficiency_benchmarks` tool allows you to compare your results against regional standards.

---

**05 How do I see if my performance is changing?**

You can use the `analyze_productivity_trends` tool to evaluate how your efficiency changes across multiple sessions.

---

**06 How do I calculate my current productivity?**

You can use the `'get_output_rate'` tool by providing the number of completed items and the total working hours.

---

**07 Can I plan future work based on my history?**

Yes, the `'calculate_capacity_requirement'` tool allows you to estimate the time needed for a target number of items using your historical rate.

---

**08 What are the regional efficiency standards?**

You can retrieve specific benchmarks for the USA or Europe using the `'get_efficiency_benchmarks'` 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>"productivity-output-rate": {<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

# Productivity Output Rate 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 Productivity Output Rate. 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 | Productivity Output Rate MCP         |
| Server ID  | 01a0ccf8-1a25-72ff-a3c6-d51a08733ac3 |
| 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/productivity-output-rate](https://vinkius.com/en/ai-agent-connect/productivity-output-rate).