
ThreatWerk User Guide

Using ThreatWerk for Threat Modeling and Supply Chain Intelligence

July 2026

1 Getting Started

1.1 First Login

Navigate to your ThreatWerk instance URL. The first user to register becomes the platform owner with full admin privileges. Subsequent users register as regular users and can be promoted by the admin.

1. Click “Create Account”
2. Enter your name, email, and password
3. You are now logged in and can create your first threat model

1.2 Inviting Users

Manage users from the Admin Area (link in the main navigation, visible to admins only).

- Users self-register at the login page
- Admins can promote users to admin role
- Access to individual models is granted per-user or via teams

2 Projects and Models

2.1 Projects

Projects are containers for organizing related threat models. They have no permissions logic - access is controlled at the model level.

- Create a project from the dashboard
- Assign models to projects via the model settings or drag-and-drop
- A model can belong to one project or none

2.2 Project Groups

Project groups aggregate multiple projects under a single compliance umbrella. They enable CRA coverage assessment across an entire product line.

- Create a project group from the dashboard
- Add projects to a group (many-to-many: a project can belong to multiple groups)
- Project groups have their own RBAC (owner/editor/viewer)
- The CRA dashboard shows aggregated compliance posture at the group level with drill-down to individual projects and models

2.3 Creating a Threat Model

1. Click “New Model” on the dashboard
2. Enter a name and optional description
3. Optionally assign to a project
4. The model opens in the diagram editor

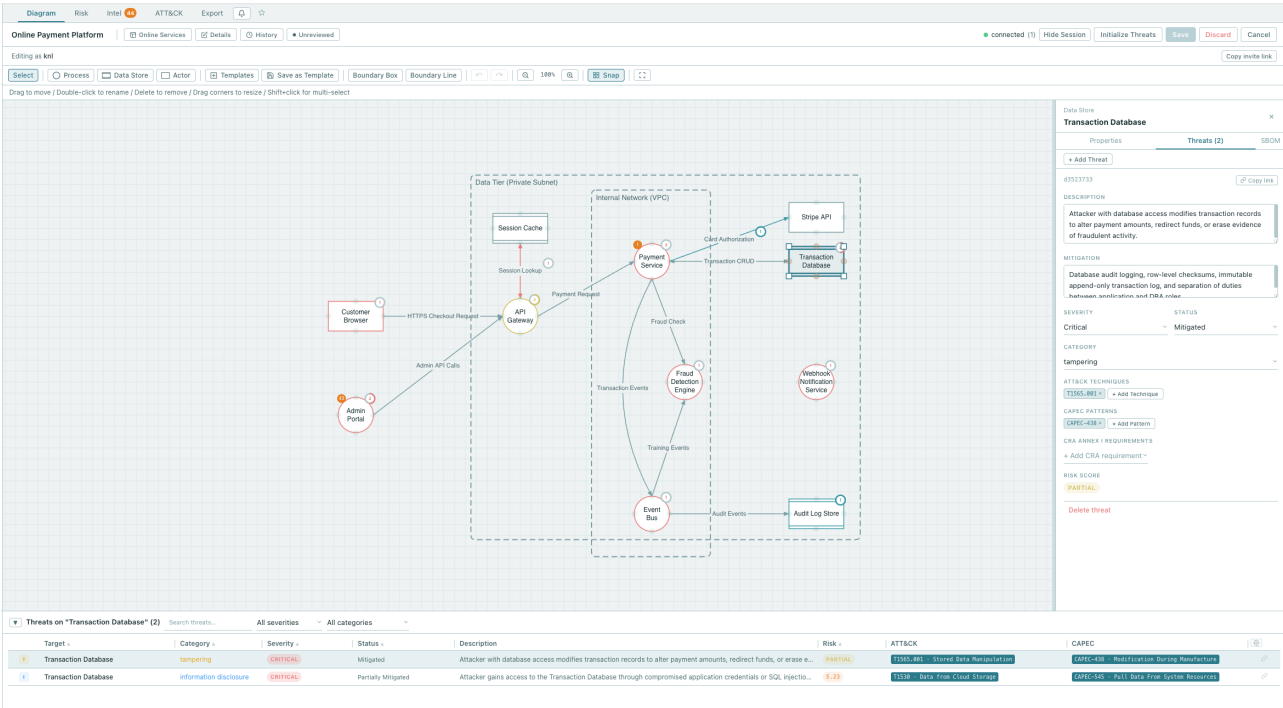


Figure 1: Diagram Editor

3 Diagram Editor

The diagram editor is the primary workspace for building threat models. It uses a canvas-based renderer for real-time collaborative editing.

3.1 Components

Components represent the building blocks of your system architecture:

Type	Visual	Use for
Process	Rounded rectangle	Services, APIs, applications, microservices
Data Store	Cylinder	Databases, caches, file systems, S3 buckets
Actor	Rectangle	Users, external systems, third-party APIs

Adding components: Click the component type buttons in the toolbar, then click on the canvas to place it. Double-click a component to edit its name.

Properties: Select a component to open the Properties Panel on the right. Here you can set: - Name and description - Component type - Supply chain tags (for intel feed matching)

3.2 Data Flows

Data flows represent communication between components.

Adding flows: Click the data flow tool in the toolbar, then click the source component followed by the destination component.

Properties: - Label (displayed on the arrow) - Protocol (e.g. HTTPS, gRPC, TCP) - Encrypted (yes/no) - Bidirectional (shows arrows in both directions) - Description

3.3 Trust Boundaries

Trust boundaries group components that share the same trust level. They can appear as dashed rectangles (area boundaries) or as dashed lines on the diagram.

Adding boundaries: Click the trust boundary tool in the toolbar. Draw a rectangular area boundary by dragging on the canvas, or draw a line boundary by clicking start and end points. Drag components into or out of area boundaries to assign membership.

Use cases: - Separate internal services from external-facing components - Mark the boundary between your VPC and the internet - Identify where data crosses security zones (critical for STRIDE analysis)

3.4 View Mode vs Edit Mode

- **View mode:** Read-only. Select components to inspect properties. Multiple users can view simultaneously.
- **Edit mode:** Full editing. Only one user edits at a time (collaborative editing shares the session). Click “Edit” to enter edit mode.

4 Threats

4.1 STRIDE Analysis

Threats are categorized using the STRIDE framework:

Category	Question
Spoofing	Can an attacker pretend to be something they are not?
Tampering	Can an attacker modify data in transit or at rest?
Repudiation	Can an attacker deny performing an action?
Information Disclosure	Can an attacker access data they should not see?
Denial of Service	Can an attacker disrupt availability?
Elevation of Privilege	Can an attacker gain higher access than intended?

4.2 Adding Threats

1. Select a component or data flow
2. Open the Threats tab in the Properties Panel
3. Click “Add Threat”
4. Select the STRIDE category
5. Describe the threat - what the attacker does and what the impact is
6. Optionally set severity, mitigation, and status

4.3 Threat Properties

Each threat has:

- **Description:** What the attacker does and the impact
- **Category:** STRIDE classification
- **Severity:** Critical, High, Medium, Low
- **Status:** Open, Mitigated, Partially Mitigated, Accepted
- **Mitigation:** How this threat is addressed
- **ATT&CK Techniques:** Mapped MITRE ATT&CK technique IDs (e.g. T1190)
- **CAPEC IDs:** Common Attack Pattern Enumeration (e.g. CAPEC-125)
- **CRA Requirements:** EU Cyber Resilience Act Annex I requirement mappings

5 OWASP Risk Scoring

ThreatWerk implements the OWASP Risk Rating Methodology for quantitative risk assessment.

5.1 How It Works

Each threat can be scored on 16 factors across 4 categories:

Likelihood Factors (8): - Skill Level (0=no technical skill, 9=security expert required) - Motive (0=no reward, 9=high reward) - Opportunity (0=no access, 9=full access) - Size (0=tiny population, 9=massive population) - Ease of Discovery (0=practically impossible, 9=trivial) - Ease of Exploit (0=practically impossible, 9=trivial) - Awareness (0=unknown vulnerability, 9=well-known) - Intrusion Detection (0=active real-time detection, 9=no detection)

Technical Impact (4): - Loss of Confidentiality (0=none, 9=all data disclosed) - Loss of Integrity (0=none, 9=all data corrupted) - Loss of Availability (0=none, 9=complete outage) - Loss of Accountability (0=fully traceable, 9=completely untraceable)

Business Impact (4): - Financial Damage (0=none, 9=bankruptcy) - Reputation Damage (0=none, 9=brand destroyed) - Non-Compliance (0=none, 9=severe regulatory penalty) - Privacy Violation (0=none, 9=millions affected)

The overall risk score and severity level are calculated automatically from these factors. Omitted factors default to worst-case (9).

5.2 Using the Risk Calculator

1. Select a threat in the Properties Panel
2. Click "Risk Score"
3. Rate each factor on the 0-9 scale
4. The overall score updates in real-time
5. Add an optional analyst comment explaining the rationale

6 Collaborative Editing

6.1 Real-Time Collaboration

Multiple users can work on the same model simultaneously:

- Changes appear in real-time for all connected users
- Cursor positions are shared between collaborators
- Conflicts are resolved automatically using operational transformation

6.2 Version History

Every change is tracked:

- View version history from the model menu
- Compare versions side-by-side (diff view)
- Each version records who made the change and when

6.3 Review Workflow

Models can go through a formal review process:

1. Author marks the model as ready for review
2. Reviewer examines the model (view mode)
3. Reviewer approves or requests changes
4. Approved models show a “Reviewed” badge

Content changes after review automatically reset the review state to prevent silent post-review tampering.

Diagram	Risk	Intel	ATT&CK	Export			
All severitiesAll statusesAll componentsSearch CVEs, tags...							
45 entries							
ID	Title	Type	Severity	CVSS	Components	ATTACK	Status
CVE-2025-44496	Axios: Regular Expression Denial of Servi...	CVE	HIGH	7.5	Admin Portal npm:axios@1.6.0	—	unreviewed
CVE-2025-44487	Axios: Proxy-Authorization Credential Le...	CVE	HIGH	8.2	Admin Portal npm:axios@1.6.0	T1552, 9861	unreviewed
CVE-2025-44486	Axios: Proxy-Authorization header leaks ...	CVE	HIGH	7.5	Admin Portal npm:axios@1.6.0	—	unreviewed
CVE-2025-44495	axios Vulnerable to Credential Theft and ...	CVE	HIGH	7.8	Admin Portal npm:axios@1.6.0	—	unreviewed
CVE-2025-44494	axios Vulnerable to Full Man-in-the-Mid...	CVE	HIGH	8.7	Admin Portal npm:axios@1.6.0	T1548, T1557	unreviewed
CVE-2025-44492	axios's shouldBypassProxy does not rec...	CVE	HIGH	8.6	Admin Portal npm:axios@1.6.0	T1552, 9861	unreviewed
CVE-2025-44490	axios has DoS & Header Injection via Pro...	CVE	MEDIUM	4.8	Admin Portal npm:axios@1.6.0	T1819, 987, T1499	unreviewed
vtx16a08021e3574a0ef2cca8a47	Copycat hits another npm package	malicious package	UNKNOWN	—	Admin Portal npm:axios@1.6.0	—	unreviewed
CVE-2025-42837	Axios: CRLF Injection in multipart/form-d...	CVE	MEDIUM	5.3	Admin Portal npm:axios@1.6.0	T1819, 987	unreviewed
CVE-2025-42838	Axios: no_proxy bypass via IP alias allow...	CVE	MEDIUM	6.8	Admin Portal npm:axios@1.6.0	T1899	unreviewed
CVE-2025-42839	Axios: unbounded recursion in toFormDa...	CVE	MEDIUM	7.5	Admin Portal npm:axios@1.6.0	—	unreviewed
CVE-2025-42834	Axios' HTTP adapter-streamed uploads ...	CVE	MEDIUM	5.3	Admin Portal npm:axios@1.6.0	—	unreviewed
CVE-2025-42836	Axios: HTTP adapter streamed response...	CVE	MEDIUM	5.3	Admin Portal npm:axios@1.6.0	T1499	unreviewed
CVE-2025-42833	Axios: Prototype Pollution Gadgets - Res...	CVE	HIGH	7.4	Admin Portal npm:axios@1.6.0	—	unreviewed
CVE-2025-42835	Axios: Header Injection via Prototype Pol...	CVE	HIGH	7.4	Admin Portal npm:axios@1.6.0	T1819, 987	unreviewed
CVE-2025-42842	Axios: XSRF Token Cross-Origin Leakag...	CVE	MEDIUM	5.4	Admin Portal npm:axios@1.6.0	T1819, T1895, +1	unreviewed
CVE-2025-42841	Axios: Authentication Bypass via Prototy...	CVE	MEDIUM	4.8	Admin Portal npm:axios@1.6.0	T1199, T1548	unreviewed
CVE-2025-42843	Axios: Incomplete Fix for CVE-2025-627...	CVE	HIGH	7.2	Admin Portal npm:axios@1.6.0	—	unreviewed
CVE-2025-42844	Axios: Invisible JSON Response Tamperi...	CVE	MEDIUM	6.5	Admin Portal npm:axios@1.6.0	T1548	unreviewed
CVE-2025-42264	Axios has prototype pollution read-side ...	CVE	HIGH	7.4	Admin Portal npm:axios@1.6.0	T1819, T1548	unreviewed
CVE-2025-42840	Axios: Null Byte Injection via Reverse-En...	CVE	LOW	3.7	Admin Portal npm:axios@1.6.0	—	unreviewed
CVE-2025-4636	Bouncy Castle has an LDAP injection	CVE	MEDIUM	5.5	Payment Service https://www.bouncycastle.org/faq.html	—	unreviewed

Figure 2: Intel Feed

7 Intel Feed

The intel feed monitors vulnerability databases and correlates CVEs to your architecture based on supply chain tags on your components.

7.1 Configuring Feed Sources

Navigate to Admin Area > Intel Feeds to configure:

Source	What it provides	API key required
NVD (NIST)	National Vulnerability Database - CVEs with CVSS scores and CPE data	Optional (higher rate limits)
GHSA	GitHub Security Advisories for open source packages	Optional (higher rate limits)
OTX	AlienVault Open Threat Exchange - community threat intel	Yes
CISA KEV	Known Exploited Vulnerabilities catalog	No
Malicious Packages (OpenSSF)	Malicious packages across npm, PyPI, etc.	No
BSI (Germany)	German Federal Office for Information Security advisories	No
NCSC (UK)	UK National Cyber Security Centre advisories	No
Cisco Talos	Threat intelligence and vulnerability research	No
Check Point Research	Threat intelligence and vulnerability research	No

Source	What it provides	API key required
Google Threat Intelligence	Google TAG threat analysis	No
Sonatype	Supply chain security research	No
CIS Advisories	Center for Internet Security advisories	No
Kaspersky Securelist	Threat research and APT reports	No
BleepingComputer	Security news and vulnerability disclosures	No
The Register	Security news and vulnerability reporting	No
Office of Inadequate Security	Data breach reporting	No

Enable feeds individually from the admin panel. Enter API keys where required. Feeds are polled automatically (default: every 12 hours, configurable via `INTEL_POLL_INTERVAL`).

7.2 Supply Chain Tags

Tags on components tell the intel feed what to match. Format:

- `npm:express` - matches npm package “express”
- `pypi:django` - matches PyPI package “django”
- `go:github.com/gin-gonic/gin` - matches Go module
- `cpe:2.3:a:apache:httpd:*` - matches by CPE

When a CVE is published that affects a tagged dependency, it appears in the model’s intel feed view.

7.3 Campaigns

Campaigns group related CVE entries for investigation:

- Create a campaign from the Intel page
- Add CVEs manually or configure auto-match rules
- Track investigation status (active, archived)
- Campaigns surface in the ATT&CK heatmap overlay

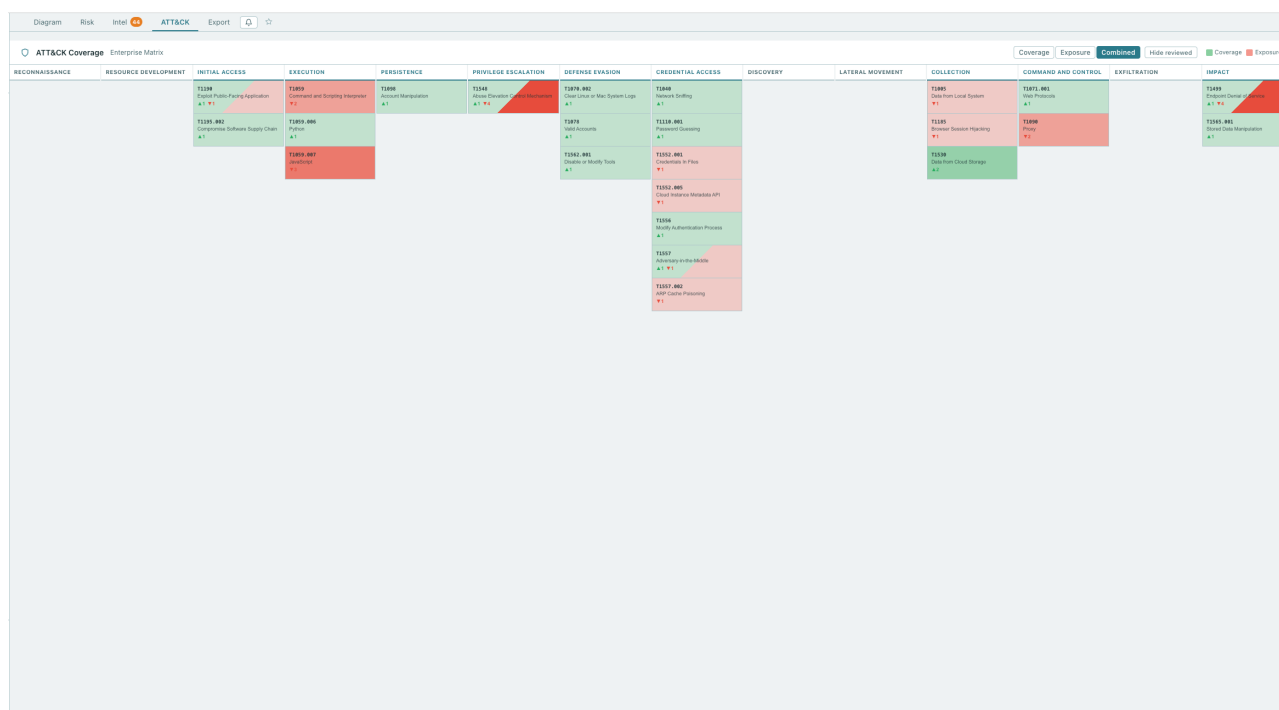


Figure 3: ATT&CK Heatmap

8 ATT&CK Heatmap

The ATT&CK heatmap overlays your threat coverage against real-world intelligence.

Coverage (blue): Techniques addressed by threats in your model. Higher coverage means more techniques are explicitly considered.

Exposure (red): Techniques associated with CVEs matched to your components via the intel feed. This shows where real vulnerabilities map to attack techniques.

Gap analysis: Techniques with exposure but no coverage represent blind spots - known attack vectors against your stack that you have not modeled threats for.

The heatmap is available per-model from the model menu.

9 SBOM Integration

SBOM (Software Bill of Materials) integration links real dependency inventories to diagram components, enabling automated supply chain tag generation.

9.1 Configuring SBOM Sources

Navigate to Admin Area > SBOM Sources:

1. Click “Add Source”
2. Configure the S3 bucket details:
 - Bucket name
 - Prefix (path within the bucket)
 - Region
 - Credentials (access key or role ARN)
3. ThreatWerk polls the bucket for SPDX JSON files

9.2 Linking SBOM Records to Components

Once records are ingested:

1. Open a model in edit mode
2. Select a component
3. In the Properties Panel, click “Link SBOM”
4. Choose the SBOM record that represents this component’s dependencies
5. Supply chain tags are automatically synced from the SBOM packages

This means intel feed matching updates automatically when your SBOM changes - no manual tag maintenance required.

10 CRA Compliance

ThreatWerk tracks coverage of the EU Cyber Resilience Act Annex I Part I requirements.

10.1 CRA Requirements

The 13 essential cybersecurity requirements:

ID	Requirement
cra_confidentiality	Protection of data confidentiality
cra_integrity	Protection of data integrity
cra_availability	Protection of availability
cra_minimize_surfaces	Minimize attack surfaces
cra_minimize_impact	Minimize impact of incidents
cra_secure_default	Secure by default configuration
cra_protect_unauthorized	Protection against unauthorized access
cra_vulnerability_handling	Vulnerability handling processes
cra_updates	Security update mechanisms
cra_monitoring	Security monitoring and logging
cra_data_minimization	Data minimization
cra_resilience	Resilience and availability under attack
cra_incident_reporting	Incident reporting capability

10.2 Mapping Threats to CRA

When creating or editing a threat, assign relevant CRA requirements. This builds a coverage matrix showing which requirements are addressed by which threats in which models.

The CRA dashboard (available from the navigation) shows overall compliance posture at two levels:

- **Project level** - aggregated across all models in a project
- **Project Group level** - aggregated across all projects in a group, with drill-down

11 Template Library

Templates are reusable architecture patterns that can be stamped onto any diagram.

11.1 Using Templates

1. In edit mode, click the Templates button in the toolbar
2. Browse or search the library
3. Click a template to preview its components and pre-defined threats
4. Click “Stamp” to place it on the diagram
5. The template’s components, data flows, trust boundaries, and threats are created with fresh IDs

11.2 Template Categories

- `aws/compute` - EC2, Lambda, ECS
- `aws/networking` - VPC, ALB, CloudFront
- `aws/storage` - S3, EFS
- `aws/database` - RDS, DynamoDB, ElastiCache
- `aws/security` - KMS, WAF, Cognito
- `aws/messaging` - SQS, SNS, EventBridge
- `generic/database` - Generic database patterns
- `generic/messaging` - Message queues, event buses
- `generic/identity` - Auth providers, SSO
- `generic/kubernetes` - K8s workloads, ingress
- `custom` - Your own patterns

11.3 Creating Templates

From existing diagram elements: 1. Select the components you want to save 2. Click “Save as Template” in the toolbar 3. Name it, assign a category, and choose scope (personal or team)

From scratch (via MCP): Use the `create_template` tool with a JSON entity definition.

12 MCP Server (AI Agent Integration)

ThreatWerk includes a built-in MCP (Model Context Protocol) server that enables AI agents (Claude, Kiro, Cursor, etc.) to read and modify threat models programmatically.

12.1 What the MCP Server Can Do

The MCP server exposes 40+ tools for:

- Listing and creating projects and models
- Adding/updating/deleting components, data flows, and trust boundaries
- Adding and scoring threats (STRIDE + OWASP risk)
- Querying the intel feed and searching CVEs
- Managing ATT&CK technique mappings
- Importing/exporting models (OTM and full format)
- Browsing and stamping templates
- Querying SBOM records and component links
- Managing campaigns

12.2 Generating an API Token

1. Click your profile avatar in the top-right
2. Go to “API Tokens”
3. Click “Generate Token”
4. Give it a name and select the scopes (read, write, admin)
5. Copy the token (prefixed with `twk_`) - it is shown only once

12.3 Configuring in Kiro / VS Code

Add to your `~/.kiro/settings/mcp.json`:

```
{
  "mcpServers": {
    "threatwerk": {
      "url": "https://your-instance.com/mcp",
      "headers": {
        "Authorization": "Bearer twk_your_personal_access_token"
      }
    }
  }
}
```

Replace `your-instance.com` with your ThreatWerk URL.

12.4 Configuring in Claude Desktop

Add to your `claude_desktop_config.json`:

```
{
  "mcpServers": {
    "threatwerk": {
      "url": "https://your-instance.com/mcp",
      "headers": {
        "Authorization": "Bearer twk_your_personal_access_token"
      }
    }
  }
}
```

12.5 Configuring in Cursor

Add to your `.cursor/mcp.json`:

```
{
  "mcpServers": {
    "threatwerk": {
      "url": "https://your-instance.com/mcp",
      "headers": {
        "Authorization": "Bearer twk_your_personal_access_token"
      }
    }
  }
}
```

The MCP server is served as Streamable HTTP at `/mcp` on the ThreatWerk backend. It uses JSON-RPC 2.0 over HTTP with Bearer token authentication. No additional bridge or proxy is required - MCP clients connect directly.

12.6 Example Workflows

Ask the AI to analyze your architecture:

“Look at my threat model ‘Payment Service’ and identify any missing threats for the data flow between the API Gateway and the database.”

Build a model from scratch:

“Create a threat model for a typical 3-tier web application with a React frontend, Node.js API, and PostgreSQL database. Add STRIDE threats for each data flow.”

Score risks:

“Score all open threats in the ‘Infrastructure’ model using OWASP risk factors. Assume a motivated external attacker with moderate skill.”

Query intel:

“Search for critical CVEs in the npm ecosystem from the last 30 days. Are any of them relevant to my models?”

12.7 Available Tools Reference

Tool	Description
whoami	Get authenticated user identity and permissions
get_version	Get server version, commit hash, and build date
list_projects	List all accessible projects
get_project	Get a project by ID
create_project	Create a new project
list_models	List threat models (optionally filter by project)
get_model	Get full threat model with components, flows, boundaries, and threats
create_model	Create a new threat model
update_model	Update model metadata
delete_model	Delete a threat model
move_model_to_project	Move model to a different project
add_component	Add a component to a model
update_component	Update component fields
delete_component	Remove a component
add_data_flow	Add a data flow between components

Tool	Description
update_data_flow	Update data flow fields
delete_data_flow	Remove a data flow
add_trust_boundary	Add a trust boundary
update_trust_boundary	Update trust boundary fields
delete_trust_boundary	Remove a trust boundary
list_threats	List threats for a model
add_threat	Add a threat (STRIDE category, ATT&CK, CAPEC, CRA annotations)
patch_threat	Update individual fields on a threat
get_risks	Get OWASP risk scores for all threats in a model
set_risk_score	Set OWASP risk score (16 factors) for a threat
get_attck_heatmap	Get ATT&CK heatmap (coverage and exposure per technique)
get_model_intel_feed	Get CVEs matched to a model's components
search_cves	Search the CVE/intel database with filters
list_campaigns	List intel campaigns
get_campaign	Get campaign details with linked entries and rules
list_teams	List teams
get_team	Get team details with members
list_model_versions	List version history for a model
list_sbom_sources	List SBOM ingestion sources
list_sbom_records	List SBOM records for a source
get_sbom_record	Get full SBOM record with parsed packages
get_component_sbom_link	Check if a component has an SBOM record linked
export_model	Export a model in OTM format
export_model_full	Export with full ThreatWerk enrichment data
import_model	Import a model from ThreatWerk full-export JSON
list_templates	Search the template library
stamp_template	Place a template onto a diagram
create_template	Create a new reusable template from entity definitions
create_template_from_selection	Save existing diagram entities as a template

Total: 42 tools covering the full threat modeling lifecycle.

13 Administration

13.1 Teams

Teams group users for shared access to models:

1. Navigate to Admin Area > Teams
2. Create a team and add members
3. Share models with teams instead of individual users
4. Team members inherit the access level granted to the team

13.2 Access Control

Model access is granted at three levels:

Role	Permissions
Viewer	Read model, view threats and intel
Editor	All viewer permissions + edit model content
Owner	All editor permissions + manage access, delete model

Access can be granted to individual users or teams.

13.3 SSO (Single Sign-On)

ThreatWerk supports OIDC-based SSO:

1. Navigate to Admin Area > SSO
2. Configure your identity provider:
 - Issuer URL
 - Client ID
 - Client Secret
3. Save and test the connection
4. Users can now log in via “Sign in with SSO”

Supported providers: Okta, Azure AD, Google Workspace, Keycloak, any OIDC-compliant IdP.

13.4 Audit Log

All significant actions are logged:

- Model creation, deletion, access changes
- User role changes
- SSO configuration changes
- API token creation
- Admin operations

Access the audit log from Admin Area > Audit Log.

13.5 Email Notifications

Configure email notifications for:

- Model shared with you
- Review requested
- Intel feed alerts (critical CVEs matching your components)

Requires `NOTIFY_BACKEND=ses` or `NOTIFY_BACKEND=smtp` and corresponding configuration. See the deployment guide for details.

14 Keyboard Shortcuts

Shortcut	Action
Delete / Backspace	Delete selected element
Escape	Deselect / cancel current tool
Ctrl+Z / Cmd+Z	Undo
Ctrl+Shift+Z / Cmd+Shift+Z	Redo
Ctrl+S / Cmd+S	Save
Mouse wheel	Zoom
Middle-click drag	Pan canvas

15 Export and Import

15.1 Export Formats

- **Image (PNG / SVG):** Diagram export for embedding in documents or presentations
- **PDF:** Printable report with diagram, threat table, and risk summary
- **OTM (Open Threat Model):** Industry-standard format for interoperability
- **ThreatWerk Full:** Proprietary format including ATT&CK mappings, CAPEC IDs, OWASP risk scores, and layout coordinates

15.2 Importing Models

- **OTM import:** Upload an OTM JSON file to create a new model
- **ThreatWerk Full import:** Restore a previously exported model with all enrichment data
- **MCP import:** Use the `import_model` tool to create models programmatically

All imports create a new model - they never overwrite existing data.