



AI your team can use.
Control your business can trust.

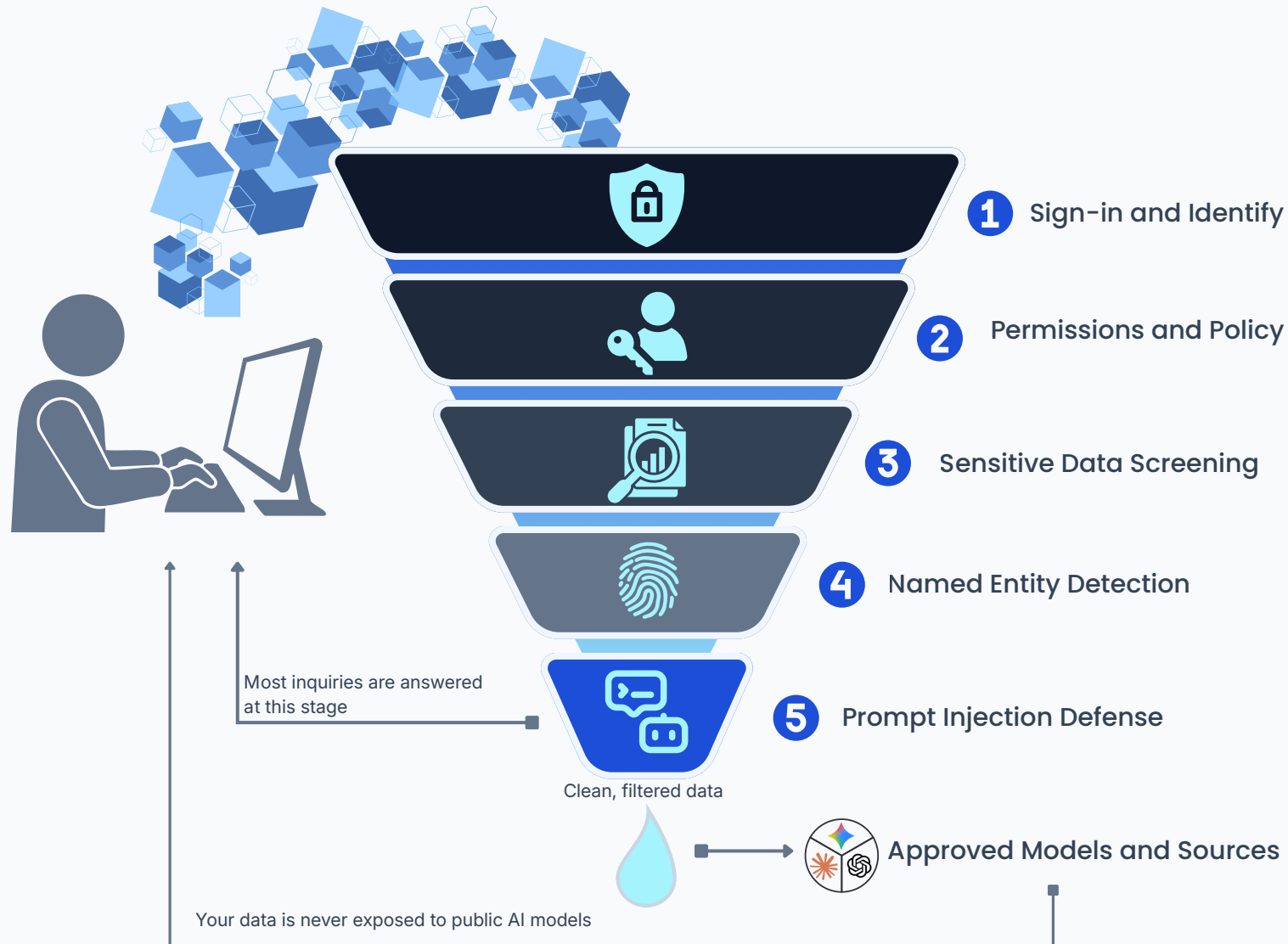
The AI Platform Built for Business

info@secure-chat.com

www.secure-chat.com

How Secure Chat Works

Every chat in Secure Chat passes through multiple layers of filtering before any public AI model is involved. Those layers are designed to stop sensitive data, block unwanted instructions, and keep routine work inside the platform, so only approved, policy-aligned information is ever sent out.



The Filtering Process

Filter #1: Sign-in and Identify

Employees sign into your company account with their Microsoft ID and multi-factor authentication. There are no personal logins, no shared passwords, and access ends the day employment does.

This prevents unauthorized access from stolen passwords or former employees accessing your account.



Filter #2: Permissions and Policies

You set access to features, models, and company data at the company level or down to the individual user, giving you precise control over how Secure Chat is used.

This prevents staff from gaining access to tools and records their role does not cover.



Filter #3: Sensitive Data Screening

Every message is checked before it moves further down the filtering process. Card numbers, account numbers, and the terms you flag are caught and held back. You decide what counts as sensitive, and can see what gets blocked.

This prevents an employee from pasting confidential information into an AI tool.



Filter #4: Named Entity Detection

Keyword lists catch only what has already been written into the rules. This layer reads the message itself, identifying people, companies, addresses, and account details, even in formats no one anticipated.

This prevents client names and personal details from slipping through because nobody listed them.



Filter #5: Prompt Injection Defense

A document, web page, or email can carry instructions to trick an AI into ignoring your rules and doing something it shouldn't. This filter treats everything arriving from outside as information to read, never as orders to follow.

This prevents a booby-trapped document quietly turning your own AI assistant against you.



Most Inquiries are Resolved Before Reaching the Public AI Models

Most routine questions, drafting tasks, and day-to-day back-and-forth are handled within Secure Chat by a private model running inside the platform. In most cases, requests are resolved before they ever need to reach an external large language model.

Public-model use is the exception, not the default. When a request does move beyond the private layer, it has already passed through the controls and protections above it. Ordinary daily conversations never leave your environment.



Approved Models and Sources

When a request genuinely needs a public model or current external information, **Secure Chat sends only the cleaned, filtered data required for that task.** Sensitive, private, and restricted information is removed or retained inside the platform before anything is shared externally. Any outside request goes only to approved providers over private business channels rather than consumer apps, and the response can show where the information came from.



With Secure Chat, you get:

- **A private workspace** for your company alone
- **End-to-end logging** across every step
- **Compliance-ready evidence** for audits and insurance reviews
- **Controls aligned to major frameworks**
- **Encryption in transit and at rest**
- **Everyday chat that stays inside the platform**
- **One admin screen** for controls and logs