



The one-page AI use *policy*

Three hundred words your team will actually read.

TIME NEEDED

An afternoon
to adopt it

YOU LEAVE WITH

A policy in force,
not in a handbook

FROM

bosio.digital
AI consulting

Your AI policy is too long. Nobody has read it.

A policy is not enforced by a document. It is enforced by people deciding, under pressure, whether the honest path is worth taking.

Most templates run to eight or twelve pages because they are written for a company with a compliance function, and because length reads as seriousness. At sixty people it buys the opposite. Every page you add reduces the chance anyone reads the four things that actually matter.

What follows is roughly three hundred words, which is three or four minutes of reading, and a document somebody can hold in their head while they work. Copy it, change the bracketed parts, delete anything that does not apply. The reasoning sits beside each clause, because you should not adopt a rule you cannot explain. The first time somebody challenges it, you will fold.

FIVE PARTS, AND THAT IS THE WHOLE POLICY

Purpose. Why this exists and who it applies to. **Tools.** What is approved, on which accounts, and what to do about anything else. **Data.** What may go in and what never may. **Release.** Who may send AI-assisted work outward. **When it goes wrong.** What actually happens, so reporting stays safe.

READ PAGE 6 BEFORE YOU SEND IT

The moment you publish this, some portion of your team is already in breach of it. Not hypothetically. That page is the part no template includes, and skipping it is how a policy converts visible risk into invisible risk.

Purpose, and tools.

Purpose

We use AI at [COMPANY] because it makes good work faster. This policy exists so that everyone knows where the edges are, and so that nobody has to guess. It applies to everyone: employees, contractors, and anyone doing work in our name.

Most policies open defensively, in the register of risk. That tells your reader the document is about protecting the company from them. Opening with why you use AI at all sets a different frame, and the boundaries land differently after that sentence than before it. Same rules, different reception.

Tools

Use [APPROVED TOOLS] on your [COMPANY] account. Do not use a personal account for company work, even on the same tool, because the company account is what keeps our data inside our agreements. If you want to use something that is not on this list, ask [NAME]. You will get an answer within [TIMEFRAME], and the answer is often yes.

This is the most important line in the policy and the one most templates omit. A sanctioned list with no route for anything else expires the moment something new appears, which in this field is roughly monthly. Worse, it teaches people that the honest path is a dead end, and the moment they learn that you lose visibility into everything. The timeframe is what makes it credible: "ask NAME" with no commitment attached is an invitation to a queue. Put a real number in and then meet it. "The answer is often yes" is doing deliberate work. It tells people the route is not decorative.

Data, release, and going wrong.

Data

These never go into any AI tool: [CLIENT CONTRACTS], [CANDIDATE AND PERSONNEL FILES], [ANYTHING UNDER NDA], [PRE-ANNOUNCEMENT FINANCIALS], [CREDENTIALS AND KEYS]. Everything else is fine. If you are unsure, the answer is to ask [NAME] rather than to decide alone, and asking is never held against you.

"Confidential information" means whatever the reader wants it to mean at four in the afternoon. Named categories do not have that problem. The other half matters just as much: everything else is fine. A list of prohibitions with no stated permission implies everything is risky, which produces exactly the hesitation you were trying to remove.

Release

You can send AI-assisted work out yourself when a mistake would be cheap and easy to undo. Anything that reaches a client, a candidate, a regulator, or the public gets read by a second person first. You are responsible for what you send, whether or not AI helped write it. "The AI wrote it" is not a defense we can offer a client, so it is not one we will offer each other.

The instinct is to gate by job title, and that fails in both directions: it slows senior people on trivial work and does nothing about a junior person sending something consequential, which is the actual risk. Sorting by what happens if it is wrong is the version that holds.

When it goes wrong

Tell [NAME] as soon as you notice. Nobody is disciplined for reporting a mistake quickly, and that is a commitment, not a comfort. What we cannot fix is the thing we find out about in six weeks from a client. Deliberately hiding a problem is a different matter, and that is treated as such.

This clause determines whether you find out about problems at all. Make a fast report punitive and you do not get fewer incidents, you get the same number and no reports, which is strictly worse. The distinction between an honest mistake and deliberate concealment has to be in the text, or the clause reads as blanket amnesty and leadership will not sign it.

Fill these in and the policy is done.

Approved tools, on company accounts

The named person who answers tool requests, and who hears about mistakes

The timeframe you will actually meet for a tool request

The categories that never go into any AI tool. Name them the way your team names them

What counts as client-facing here, so the second reader rule is unambiguous

ONE NAME, AND ONE NUMBER

If you only get two of these right, make them the named person and the timeframe. The policy's credibility is decided by the first tool request, not by the document. If the first person to ask waits two weeks, everybody watching learns what the route is actually worth.

Some of your team is already in breach. *Right now.*

Somebody has a personal account with three months of client work in it. Somebody has been pasting candidate CVs into a summarizer. Somebody built a genuinely useful thing on a tool that is not going to make the approved list.

Every template ignores this, because a template is written as though the company begins at the moment the policy is adopted. Yours does not. You have two options and only one of them works.

THE VERSION THAT FAILS

Publish the policy and let people quietly reconcile themselves to it. The existing usage does not stop. It goes further underground, because now it is a violation rather than an unaddressed habit. You have converted visible risk into invisible risk and called it governance.

THE VERSION THAT WORKS: A STATED AMNESTY, WITH A DEADLINE

Say it explicitly when you send the policy: for the next two weeks, tell us what you are using and what has gone into it, and nothing happens to you. After that, the policy applies normally.

That paragraph does three things at once. It gets you an inventory of real tool usage that no audit would have surfaced. It tells you which data has already left, which you would otherwise learn from a client. And it establishes, at the moment of introduction, that reporting is safe, which is the norm the whole policy depends on.

Expect at least one tool nobody in leadership knew about, and one category of data somewhere it should not be. That is the policy working before it has taken effect.

ONE CAUTION

Do not run an amnesty you are not prepared to honor. If somebody reports something serious and is disciplined for it, you have taught the whole company that the reporting clause is decorative, and you will not get another honest answer for years.

Writing it is an afternoon. Making it real takes a week.

Send it as a document, not a link in a handbook

Something in a handbook has been formally distributed and functionally hidden. Send it in the body of the message, and ask people to reply with one question or one case it does not cover. The questions tell you which clause is ambiguous.

Answer the first new-tool request inside your stated timeframe

The document does not decide the policy's credibility. The first request does.

Say the release rule out loud in a meeting where it applies

People adopt norms from watching them enforced once, far more than from reading them. The first time somebody says "that is client-facing, can you give it a read," the rule becomes real.

Revisit it when something surprises you, not on a calendar

An annual review produces a document that is eleven months stale for most of the year. A near miss produces a specific, correct revision. Let the edge cases write the next version.

Three signals that it is working

Adoption of the document is not the metric. Everybody clicked acknowledge, and that tells you nothing.

① New-tool requests are arriving

Counterintuitive, because a request looks like friction. It is the opposite. Zero requests over a quarter, in a team actively using AI, means people have concluded that asking is pointless.

② Near misses get reported without drama

Somebody says "I nearly pasted the wrong thing" in a normal tone, and nothing bad happens to them.

③ The awkward cases reach the named person

The test is not whether people follow the clear rules. It is what happens at the edges the policy did not foresee.

Four failures, each with one cause.

Published, and nothing changes

Almost always the tools clause. There is a list but no real route, so the first person who needs something new finds the honest path closed.

Everyone asks about everything

The data clause is written in principles rather than categories. When people cannot tell on their own, they either ask constantly or stop asking, and the second group is larger.

The second reader becomes a rubber stamp

Too much is landing in the high tier, so review became a formality. Tighten what reaches a client and the rest get real attention.

Nobody reports anything and everything seems fine

The most dangerous one, because it reads as success. Zero near misses in six months across a team actively using AI most likely means reporting does not feel safe.

WHAT WE LEFT OUT, SO YOU CAN STOP WORRYING ABOUT IT

Model-specific rules: your team chooses tools, not models. A vendor diligence process: the tools clause handles it at your size. An incident escalation matrix: if you have a security function it belongs there, and if you do not, an invented one will not be followed during a real incident. IP assignment language: your employment agreements already cover it. A prohibition on unethical use: it sounds essential and governs nothing. Grow into them if you need them.

Adopt it this week, not this quarter.

Today

Fill in the five blanks on page 5. Twenty minutes. The two that decide everything are the named person and the timeframe.

Before you send it

Decide whether you are running the amnesty on page 6, and whether you are prepared to honor it. If you are not, do not offer it.

In one quarter

Check the three signals on page 7. If all three are absent, the policy is not being followed regardless of what the acknowledgment log says, and the likeliest cause is the timeframe.

European clients or staff: the EU AI Act became generally applicable on 2 August 2026, and reaches deployers in a third country where the output is used in the Union. That is about where your work lands, not where you are incorporated.

Other guides in this set: the AI Budget Worksheet, Where to Put Your First AI Agent, The Context Files, and Twelve Questions to Ask an AI Consulting Firm.

IF YOU WANT A READ ON YOUR VERSION BEFORE IT GOES OUT

Send it over, or take 30 minutes. If we are not the right fit, we will say so and point you somewhere better.

bosio.digital/contact

Sources: the policy text, the clause reasoning, the amnesty and the three signals are published in full and free at bosio.digital, in "Your AI Policy Is Too Long. Nobody Has Read It." and the AI governance framework. EU AI Act dates from the Regulation's own text.