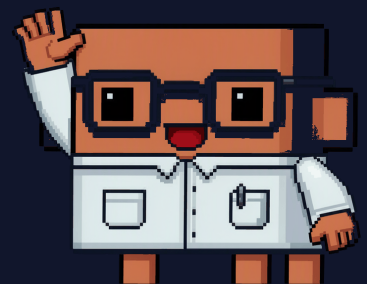


BÓVEDA · GRATIS

Claude *de cero*, en una tarde

La guía de referencia de la cuenta. Empiezas en una tarde con los cuatro primeros apartados y vuelves al resto cuando lo necesites: Proyectos, memoria, artefactos, Skills, MCP y Claude Code. Todo comprobado en las páginas oficiales, con la ruta exacta de cada ajuste.



claudeprofessor.page

Comprobado el 15 de agosto de 2026

@claudeprofessor

ÍNDICE

Los apartados son independientes. Si vienes de cero, haz **Empezar** y **Qué hacer esta tarde**. Si ya la usas a diario, salta a Skills y MCP: ahí está lo que casi nadie usa.

Comprobado el 15 de agosto de 2026 en las páginas oficiales de Anthropic. Son nombres de interfaz y se renombran sin avisar: si un menú no se llama así cuando lo abras, búscalos por lo que hace.

01 Empezar • págs. 3 a 6

Qué es Claude, en qué se diferencia, cómo pedirle las cosas y las cuatro frases que cambian el resultado.

02 Dónde vive tu trabajo • págs. 7 a 10

Proyectos, memoria, modo incógnito y los ajustes de privacidad que casi nadie toca.

03 Lo que Claude fabrica • págs. 11 a 15

Artefactos, archivos, documentos e imágenes, búsqueda web e Investigación.

04 Skills • págs. 16 a 20

Qué son, cómo se encienden, cómo se instalan en Claude Code y cómo hacerte una.

05 MCP • págs. 21 a 23

Qué es un servidor MCP, cómo se conecta y qué das cuando le das acceso.

06 Claude Code y honestidad • págs. 24 a 29

Qué es Claude Code, qué hace bien, qué hace mal y cuándo no usarla.

01 · QUÉ ES CLAUDE

Claude es un asistente de IA de Anthropic con el que se habla escribiendo. Hasta ahí, igual que las demás. Lo que cambia es otra cosa, y no es la que se cuenta en los anuncios.

La idea que ahorra más tiempo

Si te quedas con una sola frase de esta guía, que sea esta: **no es una bola de cristal, es un becario brillante que acaba de entrar.** Sabe muchísimo de casi todo, escribe rápido y no se cansa. Pero no sabe nada de ti, ni de tu empresa, ni de tus clientes, ni de lo que tú considerarías un buen resultado. Nadie se lo ha contado.



Todo lo que viene después sale de ahí. Los Proyectos existen para no volver a explicarle tu contexto cada lunes. La memoria existe para que retenga lo que le cuentas. Las Skills existen para que aprenda tu método. Y los errores del final del documento son casi todos la misma equivocación: tratarlo como un oráculo que ya debería saber, en vez de como alguien nuevo al que hay que poner al día.

Dónde se usa y qué hace falta

Funciona en la web, en las aplicaciones de escritorio de Mac y Windows, y en el móvil. Es la misma cuenta en los tres sitios y las conversaciones te siguen. Hay un plan gratuito y varios de pago —**Pro, Max, Team y Enterprise**— y la mayoría de lo que se explica aquí funciona desde el gratuito. Cuando algo necesita plan de pago, lo digo en su apartado; no vas a encontrarte con la sorpresa al final.

Los modelos, sin marear

Hay varios y se pueden elegir a mano, pero se vive perfectamente sin tocarlo. Muy resumido: **Opus 5** para lo más complejo, **Sonnet 5** para el día a día, **Haiku 4.5** cuando importa la velocidad y **Fable 5** para encargos muy largos que corren casi solos. Lo que sí conviene saber es el tamaño del contexto: los grandes leen hasta **un millón de tokens** de una vez, que viene a ser un libro entero sin trocearlo. Por eso subirle documentos largos es donde más partido se le saca, y por eso tiene su propio apartado en la página 14.

Una advertencia que vale para todos: cada modelo aprendió hasta una fecha y no sabe nada posterior. Lo peor no es que no lo sepa, es que **si no busca, no te avisa**. Rellena el hueco con lo más probable y te lo cuenta con el mismo aplomo que lo que sí sabe.

01 · EN QUÉ SE DIFERENCIA

«¿En qué se diferencia de las demás?» es la pregunta que más me hacen, y la respuesta honesta incomoda un poco: en calidad bruta se parecen mucho más de lo que dice el marketing de cualquiera de ellas. Las diferencias que sí notarás son estas cuatro, y todas son prácticas.

Trabaja sobre lo que le das

Es donde más se nota y lo que más gente me confirma después. Le subes cuarenta páginas y te contesta sobre esas cuarenta páginas, no sobre lo que se imagina que ponen. Para leer contratos, actas, pliegos e informes largos es la que menos me ha traicionado, y esa es una diferencia que solo se aprecia cuando el documento lo has escrito tú y detectas al instante si se lo está inventando.

Se le puede enseñar a trabajar

Las **Skills** son carpetas de instrucciones que Claude abre él solo cuando la conversación las pide. Escribes tu método una vez y deja de existir el «recuérdame cómo lo hago» de cada chat. Es la función más infrautilizada de todas con diferencia, está disponible desde el plan gratuito, y tiene cinco páginas en esta guía porque se explica mal en todas partes.

Se conecta a tus herramientas

El **MCP** es un estándar abierto para enchufar una IA a otros programas y datos: tu Drive, tu gestor de tareas, tu base de datos, tu calendario. Lo publicó Anthropic y hoy lo usa medio sector, así que no es una función de Claude sino algo más grande a lo que Claude llegó primero. Tiene riesgos reales que casi nadie cuenta, y por eso aquí tiene un apartado entero para él solo.

No entrena contigo por defecto

En las cuentas normales hay un interruptor que decide si tus conversaciones se usan para mejorar los modelos. No voy a decirte aquí cómo ponerlo, porque eso es cosa tuya, pero sí dónde está y qué cambia exactamente al tocarlo: está en la página 10, y cambia más de lo que la gente cree.

02 · TU PRIMER ENCARGO BIEN HECHO

La mayoría escribe una petición de una línea y juzga a Claude por la respuesta que recibe. Es el ejercicio que más rápido cambia a la gente, así que hazlo tú mismo: manda estas dos peticiones y compara.

La primera es como lo pide casi todo el mundo. No tiene destinatario, no tiene propósito y no dice cuándo estaría bien hecho. Con esa información Claude solo puede devolverte la media de todos los resúmenes posibles, que es exactamente lo que devuelve.

Resume este documento.

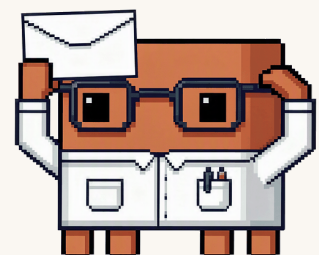
La segunda es el mismo encargo, con el mismo documento, pero ahora sabe para quién es, para qué se va a usar y con qué criterio vas a juzgarlo. Fíjate sobre todo en el final: pedirle lo que **falta** es lo que convierte un resumen en algo que sirve para decidir.

Resume esto para alguien que tiene que decidir hoy si seguimos con el proveedor. Máximo 5 puntos, y dime qué información falta para decidir bien.

La regla, en una línea

Dale un **encargo con criterio de salida**, no una pregunta. Un encargo dice para quién es, para qué sirve y cuándo está terminado. Una pregunta solo dice el tema.

Si tu petición vale para cualquiera, la respuesta también será para cualquiera. Y eso no es un defecto del modelo: es que le has dado el encargo de un desconocido.



03 · LAS CUATRO QUE SÍ CAMBIAN EL RESULTADO

De todo lo que se puede aprender sobre prompts, esto es lo que de verdad mueve la aguja. Lo demás son matices. No sustituyen a tu petición: se pegan al final de lo que ya estabas escribiendo.

Déjale preguntar

Si le falta un dato, se lo inventa antes que parar. No es malicia: es que nadie le ha dicho que puede preguntar. Esta línea sola arregla la mitad de las respuestas malas.

Pregúntame lo que te falte antes de empezar.

Pídele el diagnóstico, no el arreglo

Si le dejas reescribir, te devuelve algo correcto que suena a él y tú no aprendes nada. Pedirle qué falla te obliga a arreglarlo tú, que es como se mejora.

Dime qué frase concreta falla y por qué. No lo reescribas todavía.

Dale ejemplos de lo que quieres

Un ejemplo del resultado que buscas vale más que tres párrafos describiéndolo, y ahorra las tres rondas de «no, así no».

Te enseño un ejemplo de cómo quiero exactamente la salida.
Aplicalo a los demás sin cambiar el criterio.

Guarda lo que funcione

El error más caro es afinar algo media hora y perderlo al cerrar la pestaña. Aquí entran los Proyectos y las Skills, pero de momento basta con esto.

Conviértelo en una plantilla reutilizable con huecos claros,
y dime en qué casos NO serviría.

Un **Proyecto** es un espacio con su propio historial de conversaciones y su propia base de conocimiento. Es lo primero que deberías montar si usas Claude para algo recurrente, porque es lo que te quita de encima el repetir el contexto en cada chat nuevo.

Disponible en todos los planes. En el gratuito puedes crear **hasta cinco proyectos**. En los de pago no hay ese tope, y además el conocimiento del proyecto escala por búsqueda (RAG) en vez de ir entero al contexto, lo que multiplica hasta por diez lo que le cabe dentro.

Qué meter dentro

Documentos, textos, código o archivos que definan tu contexto estable: la guía de estilo, el pliego, el histórico de precios, las actas del año, el manual de la casa. La regla es sencilla: **no metas lo que cambia cada día, mete lo que quieres que dé por sabido siempre**. Un proyecto no es un cajón de sastre; cuanto más limpio esté, mejor contesta.

Además del conocimiento, cada proyecto lleva sus propias instrucciones. Ese es el sitio donde escribes cómo quieres que te hable, qué no debe hacer nunca y qué formato esperas. Lo que pongas ahí se aplica a todas las conversaciones del proyecto sin que tengas que repetirlo, y es lo que más se nota cuando alguien monta su primer proyecto en serio.

Cada uno recuerda por separado

La memoria de un proyecto es suya y no se mezcla con la de los demás ni con la de tus chats sueltos. Es justo lo que quieres si llevas dos clientes que no se pueden contaminar, o si usas la misma cuenta para el trabajo y para lo tuyo.

Compartirlos

En Team y Enterprise se comparten con permisos de **Puede ver** o **Puede editar**, con una persona, con varias o con toda la organización. En los planes individuales el proyecto es tuyo y solo tuyo.

Claude puede acordarse de ti entre conversaciones: a qué te dedicas, en qué andas metido, cómo te gusta que te escriban. Suena a magia y en realidad es una lista de notas que puedes leer, corregir y borrar entera.

Ojo con esto: en las cuentas nuevas la memoria es la experiencia por defecto, y la búsqueda en tus conversaciones anteriores viene activada de fábrica. Aunque no la hayas tocado nunca, merece la pena abrir esta pantalla una vez y ver qué tiene apuntado sobre ti. Se mira en el mismo sitio donde se apaga.

Dónde se enciende y se apaga

En **Ajustes → Memoria**, con el interruptor «**Generar memoria a partir de los chats**». Está en el plan gratuito y en Pro, Max, Team y Enterprise, y funciona en la web, en el escritorio y en el móvil.

Ajustes → Memoria → «Generar memoria a partir de los chats»

Al apagarla te da a elegir entre dos cosas que no son lo mismo. **Pausar** conserva lo que ya recordaba pero deja de usarlo y de añadir nada nuevo. **Restablecer** borra todas las memorias, incluidas las de los proyectos, y eso no tiene vuelta atrás. Si solo quieres dejar de alimentarla, pausa; restablecer es para empezar de cero a conciencia.

Se puede editar a mano

En la misma pantalla ves todo lo que recuerda agrupado por categorías, y hay una caja para decirle qué cambiar o qué quitar. También puedes corregirlo sobre la marcha dentro de una conversación: «eso ya no es así, apúntalo bien» funciona.

El uso que casi nadie le da es **revisarla de vez en cuando**. Si Claude te contesta raro y no sabes por qué, muchas veces la causa es una nota vieja que ya no es verdad: un cargo que cambiaste, un cliente que se fue, una manía que ya no tienes.

06 · MODO INCÓGNITO

Una conversación incógnita no se guarda en tu historial ni en la memoria de Claude. Es para lo que quieres preguntar y no quieres que quede: una duda médica, una nómina, el borrador de una carta de dimisión.

Cómo se enciende

Empieza un chat nuevo y pulsa el **icono de fantasma** de la esquina superior derecha. Sabrás que está activo porque arriba a la izquierda aparece «Chat incógnito» y la ventana se marca con un borde. Está disponible en todos los planes, del gratuito al Enterprise.

Qué no se queda

No entra en tu historial, no genera memorias, no aparece cuando buscas entre conversaciones anteriores y no cuenta para tus estadísticas de uso. Y lo más importante: **no se usa para entrenar**, aunque tengas encendido el ajuste de mejora del modelo que verás en la página siguiente.

Las dos limitaciones

La primera pilla a todo el mundo: **no funciona dentro de un proyecto**. Si empiezas la conversación dentro de uno, el fantasma no aparece por ningún lado. La segunda es que una vez cierras la conversación no se puede reabrir ni convertir en normal, así que copia lo que te interese antes de cerrarla.

Y conviene entender qué significa exactamente: por defecto se conserva **30 días** en los sistemas de Anthropic aunque tú ya no la veas. Incógnito no quiere decir que no exista; quiere decir que no se guarda contigo ni se usa para entrenar.



El ajuste que decide si tus conversaciones se usan para mejorar los modelos vive fuera de los menús que se miran a diario, y por eso casi nadie lo ha abierto nunca. Aquí está la ruta exacta y, sobre todo, qué cambia de verdad al tocarlo.

Dónde está

Pulsa tu nombre, entra en **Ajustes** y luego en **Privacidad**. El interruptor se llama «**Help Improve our AI models**» (Ayuda a mejorar nuestros modelos de IA). En el móvil la ruta es la misma, y se puede cambiar cuando quieras.

```
claude.ai/settings/data-privacy-controls
```

Lo que cambia de verdad

Aquí está el dato que casi nadie cuenta y que a mí me parece el más importante de todo el apartado: **no decide solo si entrenan contigo, decide cuánto tiempo guardan tus datos**. En las cuentas de consumo —gratis, Pro y Max—, si lo permites la retención es de **cinco años**; si no lo permites, de **30 días**. Es la misma casilla, y la diferencia entre las dos opciones no son unos meses.

Con el ajuste permitido, lo que entra es la conversación entera junto con lo que hayas adjuntado. Anthropic dice que lo desvincula de tu identidad antes de que nadie lo revise. Y hay una excepción que conviene conocer aunque no sea agradable: las conversaciones que marcan sus clasificadores de seguridad pueden usarse igualmente para sus sistemas internos de confianza y seguridad, tengas el interruptor como lo tengas.

Nada de esto lo arregla ningún ajuste: datos de clientes o pacientes, contratos con nombres, contraseñas, claves de API, DNI y números de cuenta. Si necesitas trabajar con un documento así, quítale los nombres antes. La comparativa completa con ChatGPT y Gemini está en claudeprofessor.page/privacidad.

Un **artefacto** es contenido que Claude te saca a una ventana aparte, fuera del hilo del chat: un documento, un diagrama, una web de una página, una calculadora que funciona de verdad. Aparece solo cuando lo que ha hecho es sustancial y se sostiene por sí mismo.

Requisito: hay que tener encendida la **ejecución de código y creación de archivos**, en Ajustes → Capacidades. Es el mismo interruptor que necesitan las Skills, y es donde se atasca casi todo el mundo.

Qué puede fabricar

Documentos en Markdown o texto plano, fragmentos de código, webs de una sola página en HTML, imágenes en SVG, diagramas y esquemas, y componentes interactivos en React. Todo se edita ahí mismo pidiéndole cambios, sin volver a empezar y sin que se pierda lo anterior.

Se guardan solo si los publicas

Este detalle vale por toda la página. Los artefactos que salen dentro de una conversación **no aparecen solos** en tu barra lateral: tienes que abrir el artefacto y pulsar **Publicar** para que se sume a tu colección. Si no lo haces, se quedan enterrados en el chat donde nacieron y no vuelves a encontrarlos.

Lo que casi nadie sabe que existe

Un artefacto puede llevar dentro las capacidades del propio Claude. Es decir, puedes construirte una herramienta que razone —un corrector, un clasificador, un evaluador de textos— y usarla **sin poner ninguna clave de API**. En Pro, Max, Team y Enterprise además puede guardar datos entre sesiones, hasta 20 MB, y conectarse por MCP a servicios como Asana, Google Calendar o Slack.

Ahí es donde un artefacto deja de ser una demostración bonita y empieza a ser una herramienta que usas todas las semanas.

09 · EJECUCIÓN DE CÓDIGO Y ARCHIVOS

Con este interruptor, Claude deja de describirte la tabla y te la entrega hecha. Ejecuta código en un entorno aislado y te devuelve el archivo terminado. Está en todos los planes, incluido el gratuito, y viene apagado de fábrica.

Dónde se enciende

En **Ajustes → Capacidades**, interruptor **«Ejecución de código y creación de archivos»**. En Team y Enterprise lo tiene que habilitar antes el propietario de la organización desde **Ajustes de organización → Capacidades**.

Ajustes → Capacidades → «Ejecución de código y creación de archivos»

Qué te entrega

Hojas de Excel (.xlsx) con las fórmulas vivas, presentaciones (.pptx), documentos de Word (.docx), archivos PDF, gráficos en PNG y scripts de Python. El máximo es de **30 MB por archivo**, y puedes descargarlos o mandarlos directamente a Google Drive.

La forma más rápida de entender lo que cambia es probarlo con datos tuyos. Súbele un CSV de ventas y pídele lo de abajo. La diferencia entre que te explique un análisis y que te entregue el análisis hecho es la razón por la que este interruptor abre media guía.

Con este CSV, hazme un Excel con una pestaña de resumen, una tabla dinámica por mes y un gráfico. Deja las fórmulas vivas, no valores pegados.

09 · EL AVISO QUE SÍ IMPORTA

Anthropic publica una advertencia explícita sobre esta función y merece leerse entera, porque describe el mismo riesgo que reaparece luego en los conectores y en MCP. Tiene nombre: **inyección de prompt**.

En sus propias palabras: alguien puede colar instrucciones de forma poco visible en un archivo o en una web para engañar a Claude y conseguir que **descargue y ejecute código no fiable** en el entorno aislado, o que **lea datos sensibles de una fuente conectada y los saque fuera** mediante una petición de red.

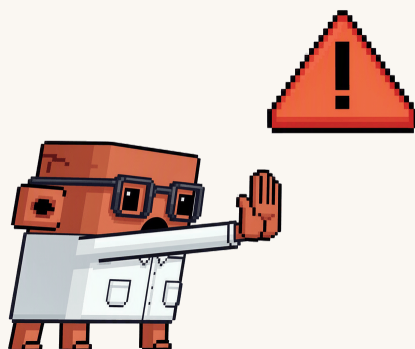
Por qué pasa esto

Claude lee lo que le llega como información, y un texto cualquiera puede llevar órdenes escondidas dentro. No separa del todo «esto es un dato que tengo que analizar» de «esto es una instrucción que tengo que obedecer» cuando ambas cosas vienen de fuera y en el mismo sitio. No es un fallo que se vaya a parchear la semana que viene: es cómo funcionan estos sistemas hoy.

Qué recomienda Anthropic

Vigilar a Claude mientras usa la función y **pararlo si ves que accede a algo que no venía a cuento**. Suena poco satisfactorio, y lo es, pero ahora mismo es la única defensa real cuando le das a la vez acceso a datos y salida a internet.

La regla práctica que yo uso: cuidado cuando coinciden **tres** cosas a la vez —contenido de fuera que no controlas, datos tuyos conectados y salida a la red—. Con dos de las tres no suele pasar nada. Con las tres es cuando se filtra algo.



10 · DOCUMENTOS E IMÁGENES

Subir archivos es lo que mejor hace y lo que peor se aprovecha. La gente sube un PDF y escribe «resume esto», que es justo la petición que menos partido le saca a la mejor función que tiene.

Los límites, en números

En una conversación puedes subir hasta **500 MB por archivo** y hasta **20 archivos**. Los PDF llegan a **1.000 páginas** y las imágenes hasta 8000 × 8000 píxeles. Acepta PDF, DOCX, CSV, TXT, HTML, ODT, RTF, EPUB y JSON, más JPEG, PNG, GIF y WebP en imagen; los XLSX necesitan la ejecución de código encendida. En el conocimiento de un proyecto el tope baja a 30 MB por archivo.

El detalle que casi nadie sabe

En PDF de **100 páginas o menos**, Claude analiza el texto **y también lo visual**: gráficos, fotografías, esquemas. A partir de 101 páginas y hasta 1.000, procesa **solo el texto**. Así que si tu documento vive de los gráficos —una memoria anual, un informe de resultados, un catálogo—, pártelo antes de subirlo. Es la diferencia entre que lea tu documento y que lea la mitad de tu documento sin decírtelo.

El encargo que lo cambia todo

En vez de pedir un resumen, pídele que se lo lea como si tuviera que defenderlo mañana delante de alguien que busca el fallo. Cambia por completo lo que te devuelve.

Lee esto y dime: las 5 ideas que de verdad importan, lo que se contradice entre secciones, y las 3 preguntas que me harían si tuviera que defenderlo.



Son dos cosas distintas y se confunden todo el rato. La **búsqueda web** es una consulta. La **Investigación** es Claude buscando por su cuenta, muchas veces seguidas, hasta cerrar el tema.

Búsqueda web

Está disponible en cuentas gratuitas y de pago, y **no viene encendida**: se activa desde el icono de ajustes de la caja de escribir, buscando «Búsqueda web». Cuando busca, cita las fuentes con enlace, que es lo que te permite comprobarlo sin fiarte.

Investigación (Research)

Solo en planes de pago —Pro, Max, Team y Enterprise— y necesita la búsqueda web encendida para funcionar. Se activa con el «+» de abajo a la izquierda, opción **Investigación**. A partir de ahí Claude encadena búsquedas, decide qué mirar después en función de lo que va encontrando y te devuelve un informe con sus citas.

Cuándo compensa cada una

Búsqueda web para un dato suelto: un precio, una fecha, si algo existe todavía.
Investigación cuando la respuesta necesita cinco o más consultas encadenadas y estás dispuesto a esperar unos minutos: comparar proveedores, actualizar un documento con lo que ha cambiado fuera, montar un estado de la cuestión antes de una reunión.

Y una cosa que hay que seguir haciendo con las dos: **comprobar las citas**. Que haya enlaces no significa que el enlace diga lo que el resumen afirma. Es el fallo que más caro sale, precisamente porque una respuesta con fuentes *parece* verificada.

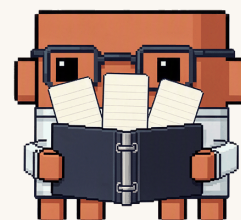
12 · SKILLS: QUÉ SON

Una **skill** no es un prompt. Es una carpeta con instrucciones —y, si quieres, con scripts y archivos de apoyo— que Claude abre él solo cuando la conversación lo pide. Escribes tu método una vez y desaparece el «recuérdame cómo trabajo» de cada chat.

Disponibles en todos los planes —gratis, Pro, Max, Team y Enterprise— con un requisito: hay que tener encendida la **ejecución de código y creación de archivos**. Sin eso no aparecen o no hacen nada.

Las que ya tienes sin saberlo

Anthropic mantiene varias skills propias para crear y manipular Excel, Word, PowerPoint y PDF. Están disponibles para todo el mundo y Claude las invoca **solo, sin que le digas nada**, cuando la tarea encaja. Si alguna vez te ha devuelto un .xlsx bien montado sin que tú mencionaras Excel, fue una skill.



Por qué no se atraganta

Funcionan por revelación progresiva: Claude mira qué skills hay disponibles, decide cuáles vienen a cuento para lo que le has pedido y carga únicamente lo que necesita. Por eso puedes tener veinte instaladas sin que se le llene el contexto ni se vuelva más torpe, que es la duda razonable que le surge a todo el mundo al oír esto por primera vez.

Qué llevan dentro

Instrucciones en Markdown, y ya está. Para una skill sencilla **no hay que programar nada**: es un documento de texto con reglas. Si quieres complicarte puedes adjuntarle scripts que se ejecuten, pero eso es el nivel avanzado y no hace falta para el 90% de los casos útiles.

Se disparan solas cuando su descripción encaja con lo que estás pidiendo, o a mano nombrándolas: «usa la skill pdf para sacar los campos de este formulario».

Si usas Claude desde la web o desde el escritorio, **no tienes que instalar nada**: las oficiales ya están ahí y solo hay que encenderlas. El problema es que son dos ajustes en dos sitios distintos, y ese es justo el paso donde se atasca todo el mundo.

Primero la ejecución de código, después las skills

El orden importa. Primero ve a **Ajustes → Capacidades** y enciende «**Ejecución de código y creación de archivos**»: las skills la necesitan para funcionar y activarlas sin eso no sirve de nada. Después entra en **Personalizar → Skills** y activa las que quieras, una a una.

Ajustes → Capacidades → «Ejecución de código y creación de archivos»
Personalizar → Skills

Si estás en un plan Team o Enterprise y no ves las opciones, no es cosa tuya: el propietario de la organización tiene que habilitarlas antes desde **Ajustes de organización → Skills**.

Subir una tuya

En **Personalizar → Skills**, botón «+», luego «**Crear skill**» y «**Subir una skill**». Le das un ZIP con la carpeta dentro. Hay dos requisitos que fallan siempre y que no avisan bien: la carpeta se tiene que llamar **igual que la skill**, y dentro tiene que ir el **SKILL.md**.

Una skill que subes queda privada de tu cuenta. Para repartirla a un equipo entero hace falta un plan Team o Enterprise, que sí permite desplegarlas a toda la organización de una vez.

Aquí sí se instalan, y hay dos caminos. El corto es copiar una carpeta. El largo, que es el que de verdad interesa, es registrar el repositorio de Anthropic como mercado de plugins y traerte las oficiales de golpe.

El camino corto

Una skill es una carpeta con un **SKILL.md** dentro. Si la copias a la carpeta de skills de tu usuario, queda disponible en todos tus proyectos. Si solo la quieres en uno, va dentro de ese repositorio.

```
~/ .claude/skills/mi-skill/SKILL.md    → en todos tus proyectos  
.claude/skills/mi-skill/SKILL.md      → solo en este proyecto
```

El mercado oficial

Primero registras el repositorio de Anthropic y después instalas los paquetes. **Son dos y necesitas los dos**, porque las skills de documentos y las de ejemplo viven separadas: si instalas solo uno te falta la mitad y acabas pensando que la skill que buscabas no existe.

```
/plugin marketplace add anthropics/skills  
/plugin install document-skills@anthropic-agent-skills  
/plugin install example-skills@anthropic-agent-skills
```

Una vez instaladas se invocan solas cuando la conversación lo pide, o a mano escribiendo **/nombre-de-la-skill**. El nombre del comando sale del nombre de la carpeta, así que si quieres cambiarlo, renombra la carpeta.

Por darte una idea de la escala: el repositorio de skills de Anthropic pasa de **169.000 estrellas** en GitHub, es de código abierto y es gratis. Trae además **skill-creator**, que sirve para escribir skills nuevas.

15 · HACERTE UNA SKILL

No hay misterio ninguno: una carpeta con un archivo dentro. El archivo empieza por un bloque con **name** y **description**, y sigue con las instrucciones que Claude va a seguir. Eso es todo lo que hay.

La descripción es la pieza crítica. Es lo único que Claude lee para decidir si la skill aplica. Si dice **cuándo** usarla con palabras concretas, se activa cuando toca. Si es vaga, o no se activa nunca o se activa siempre. Escribe esa línea como si fuera para otra persona.

Este es un SKILL.md completo y funcionando, no un esquema. La parte de arriba, entre las tres rayas, es lo que Claude lee para decidir si entra. Lo de debajo es lo que hace una vez ha entrado.

```
---
name: escribe-como-yo
description: Aprende mi voz a partir de textos míos antes
  de escribir nada. Úsala cuando te pida escribir un texto
  que tenga que sonar a mí.
---

# Instrucciones

1. No escribas nada todavía. Pide tres textos míos.
2. Descríbeme cómo escribo: longitud de frase, palabras
  que repito, qué nunca digo.
3. Espera a que confirme antes de redactar.
```

Para comprobar que funciona, abre una conversación nueva y pide algo que encaje con la descripción. Si la skill entra, verás que sigue su flujo en vez de responder directamente. Si no entra, el problema está casi siempre en la descripción, no en las instrucciones.

Dos cosas más que ahorran tiempo: saber cuándo merece la pena hacerse una, y qué campos puedes poner arriba sin que te dé error al subirla.

Cuándo hacerte una

Cuando te pillas pegando las mismas instrucciones por tercera vez. Ese es el momento exacto y no hay que darle más vueltas. Una skill no es para lo que haces una vez: es para lo que haces todas las semanas y siempre explicas igual, y por eso la primera que se hace casi todo el mundo acaba siendo la de su propio estilo de escribir.

Los campos que valen en todas partes

Las skills siguen un estándar abierto, así que el mismo archivo vale en la app, en Claude Code y en la API. Los campos que aceptan los tres son **name**, **description**, **license**, **compatibility**, **metadata** y **allowed-tools**. Si metes otros distintos, al subirla a la app te dará error, y el mensaje no siempre deja claro cuál sobra.

El atajo

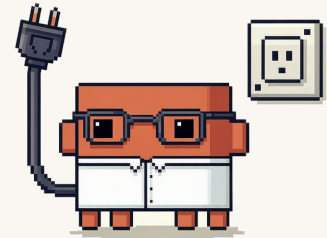
Instala **skill-creator** y pídele que la escriba por ti. Le cuentas tu método hablando y te devuelve tus propias reglas puestas por escrito, que es exactamente lo que nunca te acuerdas de repetir en cada chat.

Y si quieres las mías: tengo tres publicadas y abiertas, sacadas de métodos que he explicado en carruseles —**escribe-como-yo**, **interroga-documentos** y **reunion-a-tareas**—. Están en claudeprofessor.page/skills para descargarlas y cambiarlas a tu gusto, y el paso a paso completo está en claudeprofessor.page/instalar-skills.

MCP son las siglas de Model Context Protocol. Es un estándar abierto para conectar una IA a herramientas y datos de fuera. Lo publicó Anthropic y hoy lo usa medio sector, así que un servidor MCP no sirve solo para Claude.

La analogía que funciona

Si una skill es un manual que Claude lee, un servidor MCP son **manos**. La skill le dice cómo hacer algo; el MCP le da acceso a la cosa: tu Drive, tu gestor de tareas, tu base de datos, tu calendario. Son cosas complementarias y mucha gente las confunde porque las dos se describen como «ampliar lo que Claude puede hacer».



Lo que cambia en la práctica es que dejas de copiar y pegar. En vez de exportar el informe, pegarlo en el chat y volver a pegar la respuesta en la herramienta, le pides el resultado y él va, lo lee y lo escribe donde toca. Ahí es donde se nota que el trabajo repetitivo desaparece de verdad, y no solo se hace un poco más rápido.

Conectores y servidores

En la app de Claude se llaman **conectores**. Hay unos ya listos en un directorio, revisados por Anthropic, y luego están los **conectores personalizados**, que apuntan a un servidor MCP cualquiera. La diferencia entre esos dos grupos es todo el apartado de riesgos de la página 23, así que tenla presente desde ya.

Funcionan en Claude web, en el escritorio, en el móvil y en Claude Code. Los personalizados están disponibles desde el plan gratuito, con un tope de **uno**; en Pro, Max, Team y Enterprise puedes añadir varios.

17 · MCP: CÓMO SE CONECTA

Hay dos caminos según dónde uses Claude: en la app se hace con el ratón y en Claude Code con un comando. En los dos casos el paso importante es el mismo, y no es técnico: revisar qué permisos le estás dando.

En la app

Desde el chat, pulsa el «+» de abajo a la izquierda (o escribe la tecla «/»), pasa por **Conectores** y entra en **Gestionar conectores**. Desde ahí, otro «+» y eliges del directorio. También se llega desde los ajustes, por **Personalizar** → **Conectores**. Para uno personalizado le das un nombre y una URL, y si hace falta, las credenciales de OAuth en los ajustes avanzados.

En Claude Code

Un comando por servidor. Después, **/mcp** te dice cuáles están conectados, cuáles piden autenticación y cuáles fallan. Cuando uno sale como fallido, el problema suele ser la credencial y no el servidor.

```
claude mcp add --transport http notion https://mcp.notion.com/mcp/mcp
```

El detalle que sorprende a todo el mundo

Cuando añades un conector remoto en la app, Claude se conecta a ese servidor **desde la nube de Anthropic**, no desde tu ordenador. La consecuencia práctica es que un servidor que solo es accesible dentro de la red de tu empresa o detrás de una VPN **no va a conectar**, por mucho que tú lo veas perfectamente desde tu máquina. Es la causa número uno de «lo he configurado bien y no funciona».

Este es el apartado que no suele estar en las guías y es el que más falta hace. Conectar es fácil, lo hace cualquiera en dos minutos; entender qué acabas de conceder, bastante menos.

Lo dice Anthropic con todas las letras: al conectar un servicio le estás dando permiso a Claude para **acceder y potencialmente modificar** datos dentro de ese servicio, con tus permisos de usuario. Conecta solo servicios en los que confíes y que necesites de verdad.

Verificado no es lo mismo que personalizado

Los conectores del directorio han pasado por Anthropic. Un conector personalizado apunta a un servidor que **no ha verificado nadie**, y la recomendación oficial es clara: conectar solo con organizaciones de confianza y revisar con calma los permisos que pide antes de aceptar. Es exactamente la misma prudencia que aplicarías a una extensión del navegador, y por algún motivo aquí se aplica menos.



Es acceso, no una copia

Un conector respeta los permisos del origen: si tú no puedes ver algo en la herramienta, Claude tampoco. Pero todo lo que **sí** puedes ver, lo puede leer. Y si el conector permite escribir, puede escribir. Además, los servicios de terceros procesan los datos en su propia infraestructura y bajo sus propias condiciones, que pueden estar fuera de la Unión Europea: si manejas datos de clientes, eso no es un detalle legal menor.

Y vuelve la inyección de prompt. La documentación de Claude Code lo dice directamente: comprueba que confías en cada servidor antes de conectarlo, porque **los servidores que traen contenido de fuera te exponen a inyección de prompt**. Un servidor que lee correos o páginas web es un servidor por el que puede entrar una orden que tú no has escrito.

Claude Code es Claude trabajando dentro de tus archivos: lee el proyecto entero, edita, ejecuta comandos y hace commits. Nació para programar, pero lo interesante es que hace tiempo que dejó de vivir solo en la terminal.

Para quién es

Para quien programa, evidentemente. Pero también para quien tiene una carpeta llena de archivos y una tarea repetitiva encima: renombrar doscientos archivos con criterio, convertir un montón de documentos de un formato a otro, mantener una web sencilla. Si nunca has abierto una terminal en tu vida, hay aplicación de escritorio y versión web, así que la excusa de «eso no es para mí» se ha quedado sin base.

Dónde se usa y cómo se instala

En la terminal, como extensión de VS Code y de los IDE de JetBrains, en una aplicación de escritorio para Mac y Windows, y en el navegador en **claude.ai/code**. Es el mismo motor en todas partes, así que tus ajustes y tus servidores MCP viajan contigo de una a otra. Se instala con un comando, y luego entras en la carpeta de tu proyecto y escribes **claude**. Necesita una suscripción de pago, de Pro para arriba.

```
curl -fsSL https://claude.ai/install.sh | bash

cd tu-proyecto
claude
```

Lo que lo hace distinto es el archivo **CLAUDE.md**: un documento que pones en la raíz del proyecto y que lee al empezar cada sesión. Ahí van tus normas, tus decisiones y tus manías. Es lo mismo que las instrucciones de un Proyecto, pero para tu código.

Sin humo y por orden: esto es lo que en mi experiencia devuelve tiempo de verdad, empezando por lo que más devuelve.

Leer mucho y contestar sobre ello

Documentos largos, contratos, actas, pliegos, expedientes. Es su punto más fuerte y, curiosamente, el que menos se usa. El millón de tokens de contexto de los modelos grandes significa que le cabe un expediente entero sin trocearlo, y eso cambia la clase de preguntas que le puedes hacer: ya no es «resume», es «dónde se contradice esto consigo mismo».

Criticar sin reescribir

Pedirle qué falla y por qué, en vez de pedirle que lo arregle. Es donde más aprendes tú y donde menos suena a máquina lo que acabas publicando. Vale para tus textos y para los que te mandan a revisar.

Convertir desorden en estructura

Notas sueltas de una reunión en decisiones, tareas, responsables y fechas. Un hilo de cuarenta correos en el estado real del asunto. Un montón de respuestas de una encuesta en tres patrones. Aquí gana por goleada a hacerlo a mano, y es el uso que más rápido convence a alguien escéptico.

Explicar a tu ritmo, y fabricar el archivo

Puedes preguntarle tres veces seguidas lo que no te atreves a preguntar en una reunión, y no se cansa ni te juzga. Y con la ejecución de código encendida, la salida deja de ser texto y pasa a ser un Excel con fórmulas o una presentación montada, que es el salto que más impresiona a quien lleva años copiando y pegando.



Esta parte es la que casi ninguna guía escribe, y es la que evita disgustos. Nada de esto es un fallo puntual que se arregle con un buen prompt: son límites de cómo funciona.

Datos exactos sin fuente

Cifras, fechas, artículos de una ley, precios, referencias bibliográficas. Sin darle la fuente puede equivocarse **con total seguridad**, que es la peor combinación posible, porque no hay ninguna señal en la respuesta que te avise. Si el dato importa, dale tú el documento o enciende la búsqueda web y comprueba la cita. Relacionado con esto: tiende a rellenar antes que a parar, y por eso funciona tan bien darle permiso explícito para preguntar.

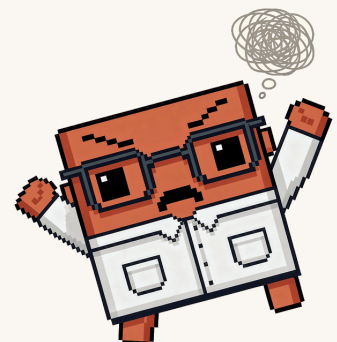
Cálculo mental

Razonar aritmética larga sin ejecutar código es de lo peor que hace. La buena noticia es que la solución está a un clic: con la ejecución de código encendida deja de estimar y calcula de verdad.

Sonar a ti

Por defecto escribe correcto y anodino. Se arregla dándole textos tuyos, no adjetivos: «hazlo más humano» no significa nada operativo, mientras que tres textos tuyos sí.

Y el límite de fondo, del que salen los avisos de las páginas 13 y 23: cuando lee algo de fuera —una web, un archivo, un servidor MCP— **no separa del todo la información de una instrucción escondida.**



Hay tareas donde el problema no es que lo haga mal. Es que usarla es una mala idea aunque lo hiciera perfectamente.

Cuando no puedes comprobar el resultado

Si no tienes forma de verificar si la respuesta es correcta y la decisión es seria, no la uses para decidir. Úsala para lo contrario: para preparar las preguntas que le vas a hacer a quien sí sabe. Eso sí lo hace muy bien y no tiene ningún riesgo.

Con datos que no puedes enseñar

Historiales médicos, datos de clientes, contratos con nombres, nóminas. Aunque tengas la privacidad bien configurada, hay información que sencillamente no debería salir de donde está. Anonimizar cuesta menos que explicar después por qué salió.

Cuando el valor está en haberlo pensado tú

Un examen, una carta de motivación, una decisión de equipo. Si el sentido de la tarea es el proceso y no el documento, delegarla es tirar el proceso a la basura y quedarte con el papel. Es la que más cara sale a largo plazo y la que menos se nota al principio.

Y una que casi nadie dice: **para lo que ya sabes hacer en dos minutos**. Hay gente que tarda diez minutos en explicarle una tarea de dos. La pregunta correcta no es «¿puede hacerlo?», sino «¿me sale a cuenta explicárselo?».

23 · ERRORES QUE VAS A COMETER

Los de siempre. Reconocerlos ahorra semanas, y los he visto en gente que lleva un año usándola todos los días.

Discutir en vez de explicar

Tres mensajes enfadado repitiendo que está mal, y ninguno diciendo qué esperabas. Enfadarse no es una instrucción. Lo mismo pasa con pegar cuarenta páginas y escribir «analiza esto» sin decir para qué: no le has dado el qué, ni el para quién.

Dejarlo todo apagado

El más caro de todos y el más fácil de arreglar. Sin la ejecución de código no hay archivos, ni artefactos, ni skills. Es **un solo interruptor** y desbloquea media guía; mucha gente lleva meses usando una versión capada de Claude sin saberlo.

Empezar de cero cada vez

Conversación suelta tras conversación suelta, repitiendo el mismo contexto cada lunes. Para eso están los Proyectos, y montar el primero cuesta cinco minutos una sola vez.

Y los dos clásicos que no pasan de moda: **no leer lo que te devuelve** —copiar y pegar sin revisar, y enterarte cuando te lo diga otro— y **usarlo solo para lo fácil**, pidiéndole textos cortos mientras sigues haciendo a mano la tarea de dos horas que repites cada semana.

Cinco cosas en orden. Las tres primeras son de diez minutos cada una, y con hacer solo la primera ya habrás sacado más partido que la mayoría.

01 Enciende el interruptor

Ajustes → Capacidades → «Ejecución de código y creación de archivos». Es lo que desbloquea archivos, artefactos y skills. Treinta segundos.

02 Mira tu privacidad

Ajustes → Privacidad, y decide tú si quieres que entrenen con tus conversaciones. Decidir es la parte importante; lo que decidas es cosa tuya.

03 Súbele un documento tuyo

Cualquier PDF que tengas pendiente de leer, de menos de cien páginas para que también vea los gráficos.

Lee esto y dime: las 5 ideas que de verdad importan, lo que se contradice entre secciones, y las 3 preguntas que me harían si tuviera que defenderlo.

04 Móntate un proyecto

Del asunto en el que más horas gastas. Métele los dos o tres documentos que siempre acabas explicando, y escribe las instrucciones del proyecto.

05 Convierte tu método en una skill

La tarea que haces cada semana igual. Pídeselo así y después súbela desde Personalizar → Skills.

Cada semana hago esto igual: [describelo]. Pregúntame lo que necesites y escíbeme un SKILL.md con name, description y las instrucciones.

Ya sabes *más que* la mayoría

En serio. La diferencia entre quien saca partido a Claude y quien no está en cuatro o cinco decisiones, no en saber trucos raros. Guarda este PDF: los apartados de Skills y MCP se entienden mejor la segunda vez, cuando ya tienes algo montado. Y lo que vaya cambiando lo voy publicando en Instagram y actualizando aquí.

@claudeprofessor

