

Cómo leer respuestas largas con NVDA

Datos de esta lección

Fase	F1 — Cimientos, seguridad y arquitectura del sistema
Módulo	M1.3 — Ingeniería de prompts adaptada a lectores de pantalla
Puntos de profesionalidad al superar el test	25

Objetivo práctico

Al terminar esta lección sabrás moverte con eficiencia por una respuesta larga de un LLM (ChatGPT, Claude o Gemini) usando encabezados, regiones, enlaces, campos y selección de texto, sin tener que escuchar todo el contenido de principio a fin cada vez.

Resultado que obtendrás

Capacidad de extraer rápidamente sólo la parte que necesitas de una respuesta larga: por ejemplo, copiar únicamente un bloque de código o un fragmento concreto, sin recorrer todo el texto.

Dónde encaja dentro del sistema final

A partir de este punto del curso, trabajarás con textos generados por IA cada vez más largos: prompts completos, documentación de proyecto, explicaciones técnicas. Esta lección es una habilidad transversal que usarás constantemente durante el resto del curso.

Conocimientos previos

- Lección «Configura el navegador para leer con comodidad con NVDA» (módulo 1.1).
- Al menos una cuenta de LLM activa (ChatGPT, Claude o Gemini).

Herramientas necesarias

- Una cuenta activa de un LLM con al menos una respuesta larga generada.

Posible coste

Esta lección no requiere ningún gasto adicional.

Modelo mental

Piensa en un informe impreso de varias páginas con un índice al principio. En lugar de leerlo de la primera a la última línea, vas directamente al apartado que necesitas gracias al índice, y dentro de ese apartado, sólo lees el párrafo relevante. Las respuestas bien estructuradas de un LLM (con encabezados, listas y bloques de código) permiten exactamente ese mismo tipo de lectura selectiva con NVDA.

Explicación

Cuando un LLM responde con un texto largo y bien formateado (con encabezados Markdown, listas y bloques de código), NVDA puede aprovechar esa estructura de la misma forma que aprovecha la estructura de una página web: saltando por encabezados con `H`, por listas con `L`, y navegando dentro de bloques de código de forma diferenciada del texto normal.

Una técnica adicional muy útil es pedirle directamente al LLM que estructure su respuesta para facilitar la lectura: por ejemplo, «Responde usando encabezados claros para cada sección» o «Pon el resultado final dentro de un bloque de código, separado de la explicación». Esto no es sólo una comodidad: es una forma legítima y recomendable de adaptar el comportamiento de la herramienta a tus necesidades de accesibilidad, en lugar de forzarte a adaptarte tú a un formato de respuesta poco cómodo.

Vocabulario nuevo

Markdown

Piénsalo así: Un sistema sencillo de anotaciones, como subrayar un título o poner viñetas a mano en un documento de papel.

Formato de texto ligero que usa símbolos sencillos (almohadillas para encabezados, guiones para listas) para dar estructura visual y semántica a un texto sin necesidad de un editor complejo.

Ejemplo: Un LLM que responde con "## Paso 1" está usando Markdown para marcar

un encabezado de nivel 2.

Preparación

Genera, en tu cuenta de ChatGPT, Claude o Gemini, una respuesta relativamente larga: por ejemplo, pide «una lista de diez preguntas frecuentes de un negocio de fontanería, con su respuesta, organizadas con un encabezado por pregunta».

Procedimiento paso a paso

1. **Contexto:** Tienes delante una respuesta larga con varios apartados.

Acción de teclado: Pulsa **H** repetidamente para saltar de encabezado en encabezado dentro de la respuesta.

Respuesta esperada de NVDA: NVDA debería anunciar cada título de sección conforme avanzas, por ejemplo «Pregunta 3, encabezado de nivel 3».

Qué significa: Puedes hacerte una idea completa de la estructura de la respuesta sin leerla entera, simplemente recorriendo sus encabezados.

Acción siguiente: Detente en el encabezado que te interese y pulsa la flecha abajo para leer el contenido de esa sección concreta.

2. **Contexto:** La respuesta incluye un bloque de código o un fragmento de texto que quieres copiar exactamente, sin errores.

Acción de teclado: Sitúa el cursor de exploración al principio del bloque, pulsa **Ctrl** + **Mayús** + **Fin** o selecciona con **Mayús** + flechas hasta el final del bloque, y pulsa **Ctrl** + **C**.

Respuesta esperada de NVDA: NVDA debería anunciar «seleccionado» o el propio texto seleccionado conforme lo marcas.

Qué significa: Has copiado exactamente el fragmento, listo para pegarlo donde lo necesites (por ejemplo, en un archivo de configuración).

Acción siguiente: Pega el contenido con **Ctrl+V** en su destino y verifica que se ha copiado completo.

3. **Contexto:** Quieres que las próximas respuestas del LLM sean más fáciles de recorrer con NVDA.

Acción de teclado: Escribe, dentro de la conversación: «A partir de ahora, estructura tus respuestas largas con encabezados claros para cada sección, y pon cualquier fragmento de código o configuración dentro de un bloque de código separado».

Respuesta esperada de NVDA: NVDA leerá la confirmación del LLM con normalidad.

Qué significa: Muchos LLM recuerdan esta instrucción durante el resto de la conversación (y, en algunos casos, se puede guardar como preferencia general de la cuenta).

Acción siguiente: Comprueba en tu siguiente pregunta si la respuesta llega mejor estructurada.

Posibles diferencias de interfaz

No todos los LLM estructuran sus respuestas de la misma manera por defecto; algunos usan más encabezados y listas de forma espontánea que otros. Si una respuesta llega como un bloque de texto corrido, sin estructura, pídele explícitamente al LLM que la reformatee con encabezados.

Errores frecuentes y diagnóstico

Una respuesta llega como un único bloque de texto sin ningún encabezado ni lista, difícil de recorrer selectivamente.

Cómo localizarlo: No es un fallo de NVDA ni del navegador: es simplemente el estilo de respuesta que ha generado el LLM en ese momento.

Solución: Pide al LLM que reformatee la misma respuesta con encabezados y listas; casi todos pueden hacerlo si se lo pides explícitamente.

Comprobación final

Genera una respuesta larga con al menos tres secciones, recórrela usando únicamente la tecla H, y localiza y copia un fragmento concreto de esa respuesta sin leerla de principio a fin. Si lo consigues, la comprobación es positiva.

Qué acabas de conseguir

Tienes una técnica de lectura selectiva que te ahorrará tiempo en el resto del curso, especialmente al trabajar con el Prompt Maestro y con la documentación de proyecto de la Fase 6.

Ejercicio práctico

Pide a tu LLM preferido una explicación de al menos cinco pasos sobre cualquier tema de este curso, exigiendo explícitamente que use un encabezado por paso, y practica saltar directamente al paso 4 sin leer los anteriores.

Resumen de lo imprescindible

- Las respuestas bien estructuradas (con encabezados, listas y bloques de código) permiten una lectura selectiva con NVDA, igual que una página web.
- Puedes pedirle explícitamente a un LLM que estructure sus respuestas para facilitar tu lectura.
- Seleccionar y copiar un fragmento concreto es más eficiente que leer toda la respuesta y luego copiar manualmente.

Estoy preparado para continuar si puedo...

- Recorrer una respuesta larga usando sólo la tecla H.
- Copiar un fragmento concreto de una respuesta sin leerla entera.
- Pedir a un LLM que estructure sus respuestas con encabezados para facilitar la lectura.

Test de comprobación

Para registrar el resultado y obtener Puntos de profesionalidad, realiza este test desde la versión web de la lección. Las opciones se muestran aquí sin señalar la respuesta correcta.

1. Recibes una respuesta muy larga de un LLM, como un único párrafo corrido sin encabezados ni listas, y necesitas encontrar sólo una parte concreta. ¿Cuál es la estrategia más eficiente según esta lección?
 - a. Leer el texto completo de principio a fin cada vez que lo necesites.
 - b. Pedirle al LLM que reformatee la respuesta con encabezados y listas, y después navegar con la tecla H.

- c. Cerrar la conversación y empezar una nueva desde cero.
2. Quieres hacerte una idea rápida de la estructura completa de una respuesta larga y bien formateada, sin leerla entera todavía. Según esta lección, ¿qué deberías hacer primero?
- a. Leer la respuesta completa con la flecha abajo, de principio a fin.
 - b. Pulsar la tecla H repetidamente para recorrer únicamente los encabezados y hacerte una idea general de las secciones.
 - c. Cerrar la respuesta y pedir al LLM que la resuma en una sola frase.
3. Un LLM responde habitualmente con un único bloque de texto corrido, sin encabezados ni listas, lo que dificulta tu lectura selectiva con NVDA. Según esta lección, ¿qué puedes hacer?
- a. Aceptar que no hay nada que hacer, porque el formato de respuesta del LLM no se puede modificar.
 - b. Pedirle explícitamente que estructure sus respuestas con encabezados claros y bloques de código diferenciados.
 - c. Cambiar de LLM cada vez que ocurra, ya que el problema es exclusivo de esa herramienta.

Fuentes

- [NVDA User Guide — NV Access \(consultado el 2026-08-30\)](#)

Última actualización de este contenido: **2026-08-30**.