

Aplicación educativa para el aprendizaje de un corpus de palabras en Náhuatl Educational application for learning a corpus of words in Nahuatl

L. A. Octaviano Bautista ^a, V. López-Morales ^a, A. Franco-Arcega ^a, A. Márquez-Grajales ^a, F. Salas-Martínez ^b,
M. A. Ojeda-Misses ^{a*}

^a Área Académica de Computación y Electrónica, Universidad Autónoma del Estado de Hidalgo, 42184, Pachuca, Hidalgo, México.

^b Área Académica de Química, Universidad Autónoma del Estado de Hidalgo, 42184, Pachuca, Hidalgo, México.

Resumen

En este artículo se presenta una aplicación educativa interactiva para el aprendizaje de un corpus de palabras de la lengua indígena Náhuatl haciendo uso de redes neuronales con memoria a largo y corto plazo (LSTM). El objetivo de la aplicación con redes neuronales es el desarrollo de un traductor de Náhuatl-Español, para preservar y promover la lengua Indígena. La aplicación ofrece una herramienta educativa accesible y moderna, utilizando tecnologías innovadoras como las redes neuronales para facilitar la comprensión y traducción de palabras y un poco de historia y cultura asociada con la lengua indígena Náhuatl. Entre el corpus de palabras a aprender están saludos, integrantes de la familia, frutas, animales y partes del cuerpo, que integran una base de datos capaz de ser traducidas en la aplicación. La aplicación permite una vez aprendidas las palabras del corpus aplicarlas y usarlas en el desarrollo de algunos juegos. La interfaz de usuario fue diseñada con el fin de ser intuitiva y atractiva, brindando desde un poco de historia hasta juegos como un memorama y un juego del ahorcado. Se aplicaron algunos instrumentos para poder evaluar (preTest) el conocimiento inicial del Lenguaje Náhuatl y (posTest) y el obtenido al final de los juegos. También se evaluaron la efectividad de la aplicación y la experiencia del usuario mediante una escala de usabilidad percibida, así como la ergonomía del sistema. Finalmente, la aplicación fue puesta a prueba con alumnos de primaria para aprender de manera constructiva e interactiva el corpus de palabras mediante los juegos mencionados.

Palabras clave: Aplicación de escritorio, aprendizaje, Náhuatl, red neuronal, LSTM

Abstract

This paper introduces an interactive educational application for learning a corpus of words from the indigenous Nahuatl language, using Long Short-Term Memory (LSTM) neural networks. The goal of the neural network-based application is to develop an educational application aimed at preserving and promoting the Indigenous language. The application is designed to provide an accessible and modern educational tool, leveraging innovative technologies such as neural networks to facilitate the understanding and translation of words, along with conveying elements of the history and culture associated with the Nahuatl language. The word corpus includes greetings, family members, fruits, animals, and body parts, all of which are integrated into a translatable database within the application. Once these words are learned, the application allows users to apply their knowledge through educational games. The user interface was designed to be intuitive and engaging, offering content ranging from historical background to games such as memory matching and hangman. Various instruments were used to assess participants' prior knowledge of the Nahuatl language (pre-test) and their knowledge after completing the games (post-test). Additionally, the effectiveness of the application and the user experience were evaluated using a usability scale, as well as the system's ergonomics. Finally, the application was tested with elementary school students, enabling them to learn the word corpus in a constructive and interactive manner through gameplay.

Keywords: Desktop application, learning, Nahuatl, neural network, LSTM.

*Autor para la correspondencia: manuel_ojeda@uaeh.edu.mx

Correo electrónico: oc45463@uaeh.edu.mx (Luis Alonso Octaviano Bautista), virgilio@uaeh.edu.mx (Virgilio López-Morales), afranco@uaeh.edu.mx (Anilu Franco-Arcega), aldo_marquez@uaeh.edu.mx (Aldo Márquez-Grajales), fernando_salas@uaeh.edu.mx (Fernando Salas-Martínez), manuel_ojeda@uaeh.edu.mx (Manuel Alejandro Ojeda-Misses).

1. Introducción

México es uno de los países con mayor diversidad lingüística en el mundo. Según la Encuesta Intercensal 2020 del Instituto Nacional de Estadística y Geografía (INEGI, 2024), más de 25 millones de personas en México se reconocen como indígenas y 7.4 millones de ellos hablan alguna lengua indígena. En los últimos años, estas lenguas indígenas han enfrentado una disminución significativa en el número de hablantes. Según el censo de población y vivienda del Instituto Nacional de Estadística y Geografía (INEGI, 2024), la población de hablantes de la lengua Náhuatl ascendía a 1.7 millones de personas, mientras que para 2020 el número disminuyó a 1.6 millones de personas (INEGI, 2024).

El descenso en el número de hablantes refleja una tendencia preocupante hacia la desaparición de estas lenguas, especialmente en comunidades donde la transmisión intergeneracional ha sido interrumpida por factores como la urbanización, la discriminación y la presión para adoptar el español como lengua principal. En el caso del Náhuatl, a pesar de contar con una base considerable de hablantes, la amenaza de la pérdida de la lengua persiste, lo que subraya la importancia de iniciativas destinadas a su preservación y revitalización (INEGI, 2024). Las lenguas indígenas son consideradas un patrimonio cultural y lingüístico invaluable de las comunidades indígenas en México y en todo el mundo. Su conservación es fundamental para preservar la identidad y la diversidad cultural, y para garantizar el acceso a la información y la comunicación intergeneracional.

El uso de la Computación y la Inteligencia Artificial (IA) con algoritmos de procesamiento de lenguaje natural ha permitido desarrollar herramientas para traducir, documentar y enseñar lenguas como el Náhuatl, zapoteco o mixteco (Santiago et al., 2024). Estas aplicaciones contribuyen a conservar el patrimonio lingüístico y cultural, así como a promover la inclusión digital de comunidades indígenas, al facilitarles acceso a tecnologías en su lengua materna y fortalecer su identidad cultural. En México, las aplicaciones digitales basadas en procesamiento de lenguaje y redes neuronales juegan un papel clave en la preservación y revitalización de lenguas indígenas en peligro de desaparecer (INEGI, 2024). La Computación y la IA han transformado la tecnología moderna, impulsando avances en automatización, procesamiento de datos y comunicación (Herrera et al., 2024).

Por su parte, las redes neuronales permiten entrenar máquinas con datos, lo que les permite reconocer patrones y estructuras del lenguaje y mejorar su precisión con el tiempo. Por ello, en este trabajo busca el desarrollo de una aplicación de escritorio interactiva para el aprendizaje de palabras en Náhuatl de la Huasteca Hidalguense. La aplicación tiene un contenido cultural e histórico, un traductor integrado basado en redes LSTM, un juego de memoria y un juego del ahorcado que emplean las palabras aprendidas por la red neuronal. La aplicación está organizada en un curso dividido en lecciones, con temas específicos, palabras y sus traducciones correspondientes al Náhuatl, así como ejemplos de oraciones que las utilizan. La red se entrena con palabras como datos que son vectorizados dada su estructura gramatical, lo que le permite codificar para traducir, reconocer patrones y generar respuestas coherentes en el idioma objetivo. La interfaz gráfica fue diseñada para ser intuitiva y atractiva para niños y niñas,

de manera que, ofrezca información histórica relevante hasta juegos interactivos.

2. Estado del arte

En la última década, diversas aplicaciones de la IA han revolucionado la traducción automática, en específico el uso de las redes neuronales aplicadas especialmente en lenguas indígenas mexicanas que enfrentan la escasez de recursos digitales y la complejidad morfológica. Donde modelos como LSTM y transformers (Kalvan *et al.*, 2021) han demostrado gran potencial para superar estos retos, mejorando la precisión y calidad de la traducción (Javed *et al.*, 2025), lo que facilita la preservación y revitalización de estos idiomas. El trabajo publicado por Hernández et al. (2024) presenta un método novedoso para la recopilación y traducción de textos del Mixteco al Español. El método consta de la recopilación, la generación de datos sintéticos con GPT-3.5, la depuración de los datos con un enfoque semiautomático, seguido de preprocesamiento y tokenización. Finalmente, es entrenado el modelo de traducción automática basados en redes neuronales recurrentes, redes neuronales recurrentes bidireccionales y transformers. En el trabajo publicado por (Bautista *et al.*, 2024) se da a conocer una propuesta para crear un traductor automático para idiomas indígenas con escasos recursos lingüísticos. Se utiliza un modelo de red neuronal transformer, el cual utiliza un corpus paralelo y algunas técnicas de alineación automática. La evaluación de la calidad de las traducciones se hizo con un corpus del idioma mixteco y se empleó una métrica de evaluación de evaluación automática BLEU (Bilingual Evaluation Understudy).

En el estudio realizado por (Zacarias, D. y Meza, I., 2021) se presenta el primer sistema de traducción automática neuronal para el idioma Ayuuk, hablado en el estado de Oaxaca por el pueblo Ayuukjä'äy (Mixes en español). El sistema propuesto se basa en la arquitectura neuronal transformer en un entorno codificador-decodificador. En los experimentos realizados se logró traducir del idioma Ayuuk al idioma Español y viceversa. El trabajo de investigación publicado por (González, 2022) da a conocer un traductor automático de oraciones simples en lengua indígena Purépecha al Español.

Para el desarrollo del modelo se utilizó una arquitectura transformer. La codificación se realizó en Google Colab y se utilizó JoeyNMT, que es un conjunto minimalista de herramientas que permite desarrollar fácilmente un modelo de traducción automático neuronal, utilizando la arquitectura transformer. En conjunto, estos trabajos evidencian que los modelos neuronales son una herramienta esencial para la traducción automática en lenguas indígenas mexicanas, permitiendo no solo mejorar la calidad de la traducción, sino también promover la inclusión digital y la preservación cultural. Por ende, en este trabajo se propone una aplicación de escritorio interactiva para el aprendizaje y traducción de palabras en Náhuatl, utilizando redes neuronales LSTM. La herramienta educativa está diseñada para preservar y promover la lengua Náhuatl del estado de Hidalgo, ofreciendo un corpus de vocabulario básico que incluye frutas y verduras, animales, partes del cuerpo y el uso de oraciones. Además, incorpora juegos como memorama y ahorcado para facilitar un aprendizaje lúdico e interactivo, con una interfaz intuitiva adaptada a estudiantes principiantes.

3. Red neuronal LSTM

3.1 Arquitectura

La aplicación de escritorio desarrollada hace uso de una red neuronal profunda llamada Long Short-Term Memory (LSTM) está basada en las redes neuronales artificiales multicapas perceptrón (MLP) y las redes neuronales recurrentes (RNN, en sus siglas en inglés) (Brownlee, 2019). La red neuronal LSTM toma una secuencia de datos como entradas y estas son propagadas a través de neuronas hacia otras de otras capas. Sin embargo, este tipo de redes contienen conexiones recurrentes para aprender de datos pasados, al igual que lo hacen las RNN (Brownlee, 2019). La diferencia entre éstas últimas, es que LSTM tiene la habilidad de aprender a superar retrasos de tiempos superiores a mil intervalos de tiempo o más, a diferencia de las RNN que suelen tener fallas en presencia entre cinco y diez intervalos de tiempo (Gers et al., 2000; Gers et al., 2002).

La red neuronal LSTM ha sido usada en gran cantidad de aplicaciones entre ellas: clasificación, generación y predicción en sistemas económicos, sociales, de control, computación, entre otras (Brownlee, 2019). Lo que permite predecir datos con mayor precisión, debido que recuerdan la información en un orden secuencial parecido a cómo lo realizan los seres humanos (Reddy & Prasad, 2018). La Fig. 1 muestra la estructura de una red LSTM, con varias unidades de recurrencia UR_n , para cada valor de la serie de tiempo. Cada UR_n contiene dos estados de memoria y tres funciones lógicas (Gers et al., 2002; Nahapetyan, 2019). Estas funciones son la puerta de entrada i_n , la puerta de olvido f_n y la puerta de salida o_n .

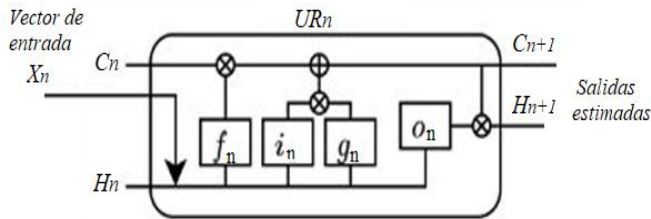


Fig. 1. Estructura de la red neuronal LSTM.

El objetivo de cada puerta es decidir cuánta información que se almacenará en cada estado, cuánto tiempo se almacenará y cuándo es necesario utilizarla dentro de la UR_n .

Por otro lado, los dos estados de memoria presentes en UR_n (son el estado de actualización C_n y el estado oculto H_n , donde el primero permite conservar la información de largo plazo, y el segundo permite conservar la información en corto plazo. La puerta de entrada i_n decide qué información del estado oculto debe ser considerada para ajustar el nivel de C_n , de acuerdo con la Ecuación 4 (Kang et al., 2020; Kang et al., 2016).

$$i_n = \sigma(W_{nx}X_n + W_{ih}H_{n-1} + b_n) \quad (1)$$

Posteriormente, para actualizar el estado C_n debe agregarse información nueva dada por la función g_n expresada en la Ecuación 5.

$$g_n = \tanh(W_{gx} + W_{gh}H_{n-1} + b_o) \quad (2)$$

La puerta de olvido f_n tiene el objetivo de decidir la información que se va a mantener dentro de la red LSTM. Para ello, toma en cuenta tanto la entrada de la célula X_n , como el estado oculto H_n devolviendo un valor entre 0 y 1, donde 0 implica que la información completa ha sido olvidada y 1 implica que la información completa se mantiene (Reddy & Prasad, 2018). La siguiente ecuación representa la estimación de f_n (Jia et al., 2017):

$$f_n = \sigma(W_{fx}X_n + W_{fh}H_{n-1} + b_f) \quad (3)$$

Finalmente, la puerta de salida o_n regula cuánta información pasará a la siguiente célula de memoria a través de H_n dada como:

$$o_n = \sigma(W_{ox}X_n + W_{oh}H_{n-1} + b_o) \quad (4)$$

Es importante mencionar que, en las ecuaciones anteriores, es necesario definir pesos para cada conexión entre funciones lógicas y estados (W_{jH}), así como entre funciones lógicas y la información de entrada (W_{jx}). Adicionalmente, es necesario agregar un conjunto de valores que definen los sesgos en cada una de las funciones lógicas (b_j).

La información de entrada de los estados de memoria de UR_{n+1} son los valores actualizados de los estados de memoria de UR_n , es decir, C_{n+1} y H_{n+1} . Las ecuaciones 5 y 6 expresan como es actualizado cada estado basado en las salidas de cada una de las funciones lógicas definidas líneas arriba:

$$C_{n+1} = (f_n C_n) + (i_n + g_n) \quad (5)$$

$$H_{n+1} = o_n * \tanh(C_n) \quad (6)$$

A continuación, se presenta el uso de la red neuronal LSTM para el procesamiento de palabras para el traductor.

3.2 Uso de la red neuronal LSTM para el traductor mediante vectorización como procesamiento de lenguaje natural

El procesamiento del lenguaje natural (PLN) ha avanzado significativamente gracias a las redes neuronales, especialmente en el ámbito de la traducción automática. Inicialmente se empleaban métodos basados en reglas y posteriormente modelos estadísticos; sin embargo, las redes neuronales mejoraron la calidad de las traducciones y permitieron una interpretación más contextual del lenguaje (Liu et al., 2020). Dentro de este contexto, las arquitecturas de redes neuronales recurrentes con memoria a largo y corto plazo (LSTM) se consolidaron como una de las más efectivas, al superar el problema del desvanecimiento del gradiente presente en las RNN tradicionales y permitir capturar

dependencias a largo plazo en secuencias lingüísticas (Sutskever et al., 2014). La combinación de LSTM con modelos sequence-to-sequence y mecanismos de atención marcó un hito en la traducción automática neuronal (NMT).

Paralelamente, la vectorización de palabras mediante técnicas como Word2Vec, GloVe, FastText y embeddings contextualizados ha permitido enriquecer la representación semántica del texto, favoreciendo el aprendizaje de correspondencias entre lenguas. Esta combinación de arquitectura (LSTM) y modelo de vectorización sigue siendo relevante, sobre todo en lenguas con recursos limitados, como las lenguas indígenas en riesgo de desaparición.

Modelos más recientes basados en la arquitectura Transformer, como BERT, GPT y LLaMa, han ampliado la capacidad de entender el contexto, mejorar tareas de traducción, clasificación y análisis de sentimientos (Devlin et al., 2019). La vectorización de palabras con Word2Vec, GloVe y BERT ayuda a capturar mejor las relaciones semánticas. Sin embargo, vectorizar idiomas indígenas como el náhuatl representa un reto debido a su complejidad morfológica, variabilidad dialectal y escasez de corpus etiquetados (Liang et al., 2020). Para el náhuatl, es fundamental contar con un corpus adecuado que incluya distintas variantes dialectales y refleje su estructura morfológica basada en sufijos y prefijos (Torres-Moreno et al., 2024; Gutiérrez-Vasques et al., 2016). Proyectos como π -YALLI y Axolotl representan bases importantes para trabajar con esta lengua en PLN (Pugh & Tyers, 2023).

3.2.1 Representación numérica de caracteres

La vectorización de palabras es un proceso crítico en el Procesamiento del Lenguaje Natural (PLN), ya que transforma palabras en representaciones numéricas, permitiendo que los modelos de aprendizaje automático, especialmente las redes neuronales, trabajen eficazmente con datos textuales. El objetivo principal es convertir palabras, que son cualitativas, en datos cuantitativos procesables por algoritmos (HaCohen-Kerner, Miller, & Yigal, 2020). Matemáticamente, el alfabeto se define como un conjunto ordenado. Si denotamos el conjunto de caracteres del alfabeto Náhuatl como A , entonces:

$$A = \{a_1, a_2, \dots, a_n\} \quad (1)$$

donde $a_1, a_2, \dots, a_n$ representa los caracteres del alfabeto Náhuatl, ordenados según su posición. En este caso, el conjunto incluye no solo los caracteres estándar del alfabeto español, sino también las vocales acentuadas y la letra "ñ". La posición de cada carácter en el conjunto se puede usar para asignarle un valor numérico real. Entonces, un carácter $c \in A$ se define por su representación como:

$$f(c) = \frac{\text{pos}(c)+1}{n} \quad (1)$$

donde $\text{pos}(c)$ es la posición del carácter c y n es el número total de caracteres en el alfabeto. Es importante mencionar que, si el carácter no pertenece al alfabeto, se le asigna el valor 0.

3.2.2 Vectorización de palabras

La vectorización de palabras es un proceso crítico en el Procesamiento del Lenguaje Natural (PLN), ya que transforma palabras en representaciones numéricas, lo que permite que los modelos de aprendizaje automático, especialmente las redes neuronales, trabajen de manera eficaz con datos textuales. El objetivo principal es convertir palabras, que son cualitativas, en datos cuantitativos procesables por algoritmos (HaCohen-Kerner, Miller, & Yigal, 2020).

La vectorización constituye un paso esencial en el PLN, dado que convierte el texto en representaciones numéricas adecuadas para el análisis computacional. Entre las técnicas más comunes se encuentran la codificación one-hot, la frecuencia de término-frecuencia inversa de documento (TF-IDF) y los métodos de representación distribuida. A pesar de su simplicidad, TF-IDF continúa siendo ampliamente utilizada debido a su efectividad en tareas como análisis de sentimientos y clasificación de textos (Kurniawan & Maharani, 2020; Nurdin et al., 2020; Çelik & Koç, 2021; Afifah, Yulita & Sarathan, 2021).

Por otro lado, enfoques más avanzados, como Word2Vec, GloVe y FastText, permiten generar embeddings densos que capturan relaciones semánticas y sintácticas entre palabras, superando las limitaciones de los modelos basados únicamente en frecuencia. Estas representaciones han demostrado mejorar significativamente el desempeño de los sistemas de traducción automática, la recuperación de información y otras tareas complejas de PLN. Por ende, debido a que las palabras son representadas en datos numéricos, una palabra se transforma en un vector de longitud fija. Supongamos que una palabra está formada por una secuencia de caracteres $X = (c_1, c_2, \dots, c_m)$. El vector de la palabra es definido como:

$$\vec{X} = (f(c_1), f(c_2), \dots, f(c_m)), \quad (3)$$

Si $m < \text{longitud}$ del vector, se rellena con ceros hasta alcanzar la longitud deseada. Si $m > \text{longitud}$, se trunca a los primeros 15 elementos. Esta normalización asegura que todos los vectores tengan la misma dimensión, lo cual es un requisito para procesarlos con la red neuronal LSTM. En este caso, la longitud de las palabras es a lo más de 15 letras, por lo tanto, el vector se define como:

$$\vec{X} = (f(c_1), f(c_2), f(c_3), 0, 0, \dots, 0) \in \mathbb{R}^{15} \quad (4)$$

Finalmente, el corpus vectorizado se construye aplicando este proceso a cada palabra de la lista de palabras en Náhuatl. Por ejemplo, si tenemos una lista de tres palabras en Náhuatl (por ejemplo: atl, calli y xochitl), cada palabra se transforma en un vector de longitud fija, y el conjunto completo de vectores se presenta como una matriz definida como:

$$C = \begin{bmatrix} \vec{X}_1 \\ \vec{X}_2 \\ \vec{X}_3 \end{bmatrix} \in \mathbb{R}^{3 \times 15} \quad (5)$$

La matriz C es el corpus de palabras de entrada, definido como un conjunto numérico que representa palabras en una forma interpretable por la red neuronal LSTM. Esta permite a la red interpretar, procesar y aprender patrones lingüísticos codificados en las representaciones vectoriales de las palabras.

Formalmente, $C \in \mathbb{R}^{k \times m}$ actúa como el tensor de entrada, donde cada fila corresponde a una palabra vectorizada de longitud fija m y k es el número total de palabras en el corpus. Esto es:

$$C = \begin{bmatrix} \vec{X}_1 \\ \vec{X}_2 \\ \vdots \\ \vec{X}_n \end{bmatrix} \in \mathbb{R}^{k \times m} \quad (6)$$

Esta entrada estructurada permite que la red identifique dependencias tanto a nivel de carácter como de palabra, lo cual es crucial para manejar lenguas morfológicamente complejas como el Náhuatl.

4. Resultados experimentales

Los resultados experimentales se dividen en dos secciones. Primero, son presentados los resultados de entrenamiento de la red neuronal para la identificación de las palabras en Náhuatl, posteriormente, el entrenamiento es usado para el diseño del traductor y en el juego del ahorcado. Finalmente, la segunda etapa consta de la integración de la aplicación y los contenidos que son descritos a continuación.

4.1 Entrenamiento de la red neuronal LSTM

La arquitectura específica del modelo de lenguaje empleado en la aplicación se basa en redes neuronales con memoria a largo y corto plazo (LSTM), seleccionadas por su capacidad para manejar secuencias de texto y capturar dependencias a corto y largo plazo entre palabras. El núcleo del traductor Náhuatl-Español está conformado por un esquema encoder-decoder. El encoder LSTM procesa las secuencias de palabras en náhuatl, convirtiéndolas en representaciones vectoriales que preservan tanto el orden como el contexto semántico. El decoder LSTM toma estas representaciones y genera la traducción al español de manera secuencial, prediciendo un token a la vez. La arquitectura incluye capas de embeddings, que transforman cada palabra del corpus en vectores numéricos, y mecanismos de regularización para evitar sobreajuste. El entrenamiento se realizó sobre un corpus de palabras en náhuatl con contenido cultural e histórico, lo que permite no solo traducciones, sino también la preservación del valor patrimonial.

Una vez dada la red neuronal se procedió a la fase de desarrollo. A continuación, se describe el proceso de generación de código para desarrollar la red neuronal LSTM en Python. Primeramente, se realizó la importación de las bibliotecas de Python Keras, biblioteca que facilita la creación de modelos para el aprendizaje profundo y construidas sobre Tensor Flow, biblioteca de código abierto desarrollada por Google que se utiliza ampliamente en el campo del aprendizaje profundo y el procesamiento de datos.

Para la implementación se utilizó la biblioteca pandas para leer y cargar los datos de los archivos llamados 'palabras_espanol.csv' y 'palabras_nahuatl.csv', los cuales contienen una columna: 'palabra' y 'traducción' respectivamente. Posteriormente se añadieron marcadores de inicio (start) y fin (end) a las traducciones para que el modelo

sepa cuándo comenzar y terminar la secuencia de salida. Para el módulo de preprocesamiento se realizó la tokenización de los datos, es decir, se convirtieron las palabras a vectores para poder ser procesados por el modelo.

El listado de palabras consta de setenta y cuatro palabras en Español y su respectiva equivalencia en Náhuatl. Una clasificación sencilla es animales: águila, ajolote, araña, borrego, búho, caballo, cerdo, conejo, cotorro, cucaracha, gallina, gallo, gato, golondrina, grillo, guajolota, gusano, hormiga, jabalí, lagartija, mapache, mariposa, mosca, paloma, pato, perro, pez, piojo, pollito, rana, ratón, sapo, serpiente, tlacuache y venado (34 palabras). Frutas y alimentos, entre ellas se tienen: la palabra fruta, chalahuite, coco, durazno, guayaba, jícama, limón, mamey, manzana, melón, naranja, piña, plátano, sandía, tejocote y zapote (16 palabras). Finalmente, algunas partes del cuerpo humano como incluida la palabra cuerpo, barba, boca, brazo, cabello, cachete, cara, carne, ceja, codo, cuello, dedo, diente, hombro, hueso, lengua, mano, nariz, ojo, oreja, panza, pie, rodilla, uña (24 palabras).

En el módulo de entrenamiento se definió la arquitectura del modelo de red neuronal LSTM utilizando Tensor Flow y Keras. El modelo consta de una capa LSTM con 256 neuronas, esto se especifica en el parámetro `latent_dim` que posteriormente se pasa al constructor de la capa LSTM. La capa de salida del modelo, que es una capa densa (Dense) utiliza la función de activación softmax para generar las probabilidades de las palabras del vocabulario de destino (Náhuatl).

El traductor funciona mostrando una ventana de salida (ver Fig. 2) cuando el error cae por debajo de un umbral predefinido, denotado como delta (δ). Esto indica que la red neuronal ha aprendido exitosamente el comportamiento deseado utilizando el algoritmo propuesto. Por el contrario, si el error permanece por encima de delta, el programa genera un mensaje notificando al usuario que el error excede el límite aceptable. Este resultado sugiere que la red neuronal aún no ha aprendido de manera óptima el patrón objetivo con el algoritmo de entrenamiento actual.

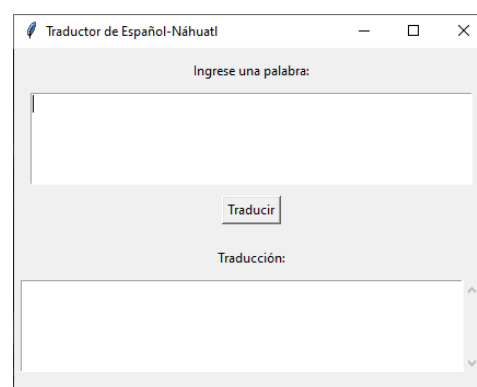


Fig. 2. Ventana del traductor que funciona mediante la red neuronal LSTM.

La parte del preprocesamiento se realiza cuando se introduce una palabra al traductor, por ejemplo, kwaitli, donde el sistema primero se procesa letra por letra. Cada carácter se busca en el alfabeto definido dentro del código (aábcdeéfgghijiklmnnñoóppqrstuúvwxyz'). Si el carácter está en el alfabeto, se convierte en un número entre 0 y 1 de acuerdo con

su posición relativa; por ejemplo, la x ocupa una posición hacia el final, por lo que su valor será cercano a 1, mientras que la o tendrá un valor intermedio. Así, la palabra se convierte en una secuencia de números que conserva indirectamente la forma de la palabra.

Luego, para garantizar que todas las palabras tengan la misma longitud de representación (en este caso, 15), el vector se completa con ceros al final hasta llegar a ese tamaño. De esta manera, kwaitli tiene 7 letras, se representará como un vector de 15 elementos: los primeros 7 corresponden a las letras convertidas en números y los 8 restantes serán ceros.

Este procedimiento asegura que tanto palabras cortas como largas tengan un formato homogéneo, lo cual es fundamental para que la red neuronal pueda comparar, entrenar y reconocer patrones lingüísticos sin importar las diferencias en la longitud de las palabras. Para la aplicación presentada se realiza la vectorización de las 74 palabras, aquí solo se presenta el ejemplo para 5 palabras: kwaitli, axolotl, tokatl, ixkatl, tekolotl.

Palabra: kwaitli [0.4375, 0.8125, 0.03125, 0.34375, 0.78125, 0.46875, 0.34375, 0.0, 0.0, 0.0, 0.0, 0.0, 0.0, 0.0, 0.0]

Palabra: axolotl [0.03125, 0.9375, 0.59375, 0.46875, 0.59375, 0.78125, 0.46875, 0.0, 0.0, 0.0, 0.0, 0.0, 0.0, 0.0, 0.0]

Palabra: tokatl [0.78125, 0.59375, 0.4375, 0.03125, 0.78125, 0.46875, 0.0, 0.0, 0.0, 0.0, 0.0, 0.0, 0.0, 0.0]

Palabra: ixkatl [0.34375, 0.9375, 0.4375, 0.03125, 0.78125, 0.46875, 0.0, 0.0, 0.0, 0.0, 0.0, 0.0, 0.0, 0.0]

Palabra: tekolotl [0.78125, 0.1875, 0.4375, 0.59375, 0.46875, 0.59375, 0.78125, 0.46875, 0.0, 0.0, 0.0, 0.0, 0.0, 0.0]

Los resultados obtenidos utilizando la red neuronal LSTM se presentan en la Fig. 3. Cabe mencionarse que en dicha figura solamente se presenta la identificación de cinco palabras vectorizadas y sus errores cuadráticos medios. Las palabras empleadas están dadas por un corpus de categorías dadas por salud, integrantes de la familia, frutas, animales y partes del cuerpo. En cuanto al error cuadrático medio (ECM) la identificación de las palabras tiene un índice bajo, lo que indica una condición inicial favorable. A lo largo del proceso de entrenamiento, el ECM disminuye rápidamente y se estabiliza en un valor en estado estable significativamente menor. En consecuencia, la red entrenada con esta función de activación logra una mejor precisión y una convergencia más rápida, lo que la convierte en una opción más adecuada para la tarea de identificación permitiendo usar la red neuronal LSTM para el traductor Náhuatl-Español aplicado en este trabajo.

La Tabla 1 muestra la evolución de los ECMs a lo largo del tiempo (de 0 a 100 unidades) para las palabras kwaitli, axolotl, tokatl, ixkatl, tekolotl. El ECM es una métrica común utilizada para medir la diferencia cuadrática promedio entre los valores predichos y los reales. Valores más bajos de ECM indican un mejor ajuste y un desempeño más preciso del modelo. La tabla revela una tendencia clara de disminución del error en todos los modelos a medida que avanza el tiempo, lo que sugiere un proceso de aprendizaje o adaptación en el que cada sistema mejora sus predicciones con el tiempo.

Los resultados de los errores cuadráticos medios (ECM) muestran una clara tendencia de convergencia hacia cero a medida que aumenta el número de iteraciones, lo que indica que la red neuronal está aprendiendo eficazmente a reproducir los vectores de referencia normalizados de cada palabra.

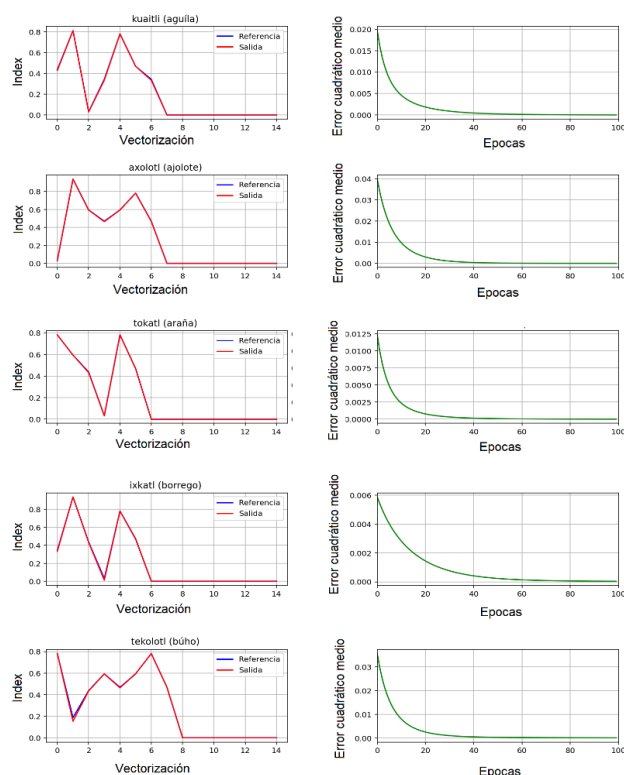


Fig. 3. Resultados obtenidos de la red neuronal mediante la vectorización de palabras y sus respectivos errores cuadráticos medios.

Por ejemplo, palabras como "kuaitli" y "axolotl" comienzan con ECM relativamente altos en la iteración inicial (0.0196 y 0.0398 respectivamente) pero disminuyen rápidamente a valores cercanos a cero en la iteración 100 (0.000018 y 0.000005), evidenciando que la función de activación y el proceso de ajuste de pesos permiten que la salida de la red se acerque progresivamente al vector esperado. Esta reducción sostenida del ECM es consistente en la mayoría de las palabras, reflejando la estabilidad y efectividad del algoritmo de aprendizaje.

Sin embargo, algunas palabras como "kuatochi" muestran una convergencia más lenta, donde el ECM en la iteración 100 todavía es ligeramente mayor (0.000174) que en otros casos. Esto podría deberse a la mayor complejidad de la palabra o a la longitud y distribución de los valores en su vector normalizado, que requieren más iteraciones o ajustes finos de los pesos para alcanzar una aproximación perfecta. En general, los resultados demuestran que el método de vectorización entera normalizada combinado con la actualización adaptativa de pesos logra una reducción significativa del error, asegurando que la red neuronal aprenda correctamente las representaciones internas de las palabras, lo que valida la efectividad del modelo para tareas de traducción entre Náhuatl y Español.

Tabla 1: Errores cuadráticos medios para la identificación de las palabras vectorizadas kuaitli, axolotl, tokatl, ixkatl, tekolotl mediante la red LSTM.

Iteración	kuaitli	axolotl	tokatl	ixkatl	tekolotl
0	0.019646	0.039857	0.012365	0.005873	0.035621
25	0.001273	0.001753	0.000487	0.001043	0.001550
50	0.000255	0.000164	0.000073	0.000233	0.000281
75	0.000062	0.000022	0.000013	0.000072	0.000124
100	0.000018	0.000005	0.000003	0.000038	0.000080

4.2. Aplicación de escritorio

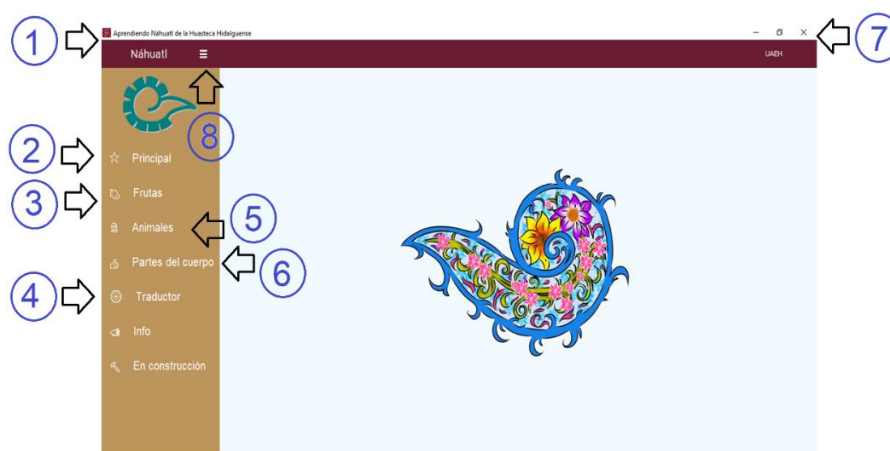


Fig. 4. Ventana principal de la aplicación de escritorio para aprendizaje de Náhuatl.

Una vez presentada la identificación de las palabras usando la red neuronal LSTM es usada para el diseño del traductor y del juego del ahorcado para el contenido de la aplicación de escritorio (ver Fig. 4).

En la aplicación de escritorio son presentados algunos elementos que son: 1. La barra de encabezado, que contiene el nombre de la aplicación de escritorio. 2. El menú principal donde se da un preámbulo sobre la historia e importancia de la lengua Náhuatl en México. 3. Un apartado donde se da el contenido de palabras referentes a las frutas, algunas oraciones formuladas usando dichas palabras y dos juegos (un memorama y un juego de ahorcado). 4. En este apartado se crea el menú de traductor donde el usuario puede traducir un corpus de palabras dadas por saludos, frutas, animales y partes del cuerpo de Náhuatl al Español y viceversa. 5. En la pestaña de los animales se tiene un listado, algunas oraciones formuladas usando dichas palabras y dos juegos (un memorama y un juego de ahorcado). 6. En la pestaña de partes del cuerpo se tiene un listado de ellas, algunas oraciones formuladas usando estas palabras y dos juegos (un memorama y un juego de ahorcado).

Finalmente, en 7. Se tiene las funciones esenciales de una ventana, es decir, dicha ventana puede minimizarse totalmente, minimizar y maximizar el tamaño y cerrar la ventana, donde cada parte es enumerada en la Fig. 4

La estructura de la navegación del contenido de la aplicación de escritorio está basada en el siguiente. Primero la página principal contiene los apartados: Aprendiendo sobre el Náhuatl, ¿Dónde se originó la lengua Náhuatl?, la relevancia del Náhuatl en el periodo colonial y las características de la lengua Náhuatl. En cuanto a las lecciones se tiene la de frutas, que tiene una presentación de lección, una tabla de frutas en Español y Náhuatl, una tabla de oraciones sobre frutas en Español y Náhuatl y los juegos.

La sección de animales contiene la presentación de la lección, la tabla de animales en Español y Náhuatl, la tabla de oraciones sobre animales en español y Náhuatl y los juegos; para la sección de partes del cuerpo se considera muy similar con el contenido: presentación de lección, la tabla de partes del cuerpo en español y Náhuatl, la tabla de oraciones sobre partes del cuerpo en español y Náhuatl y los juegos; finalmente, una pestaña para el traductor En la Fig. 5 se presenta la primera ventana del menú principal. Aquí se describe detalles de la lengua Náhuatl, un número aproximado de personas que hablan dicha lengua y algunos estados donde se habla. En el apartado de topico de frutas se presenta una lista con algunas frutas como el coco, el durazno, la guayaba, el limón, la piña, el platano y la sandía; donde ademas de contener su traducción al Náhuatl se presentan ejemplos de oraciones (ver Fig. 6).



Fig. 5. Ventana inicial de la aplicación de escritorio que describe datos de la lengua Náhuatl.

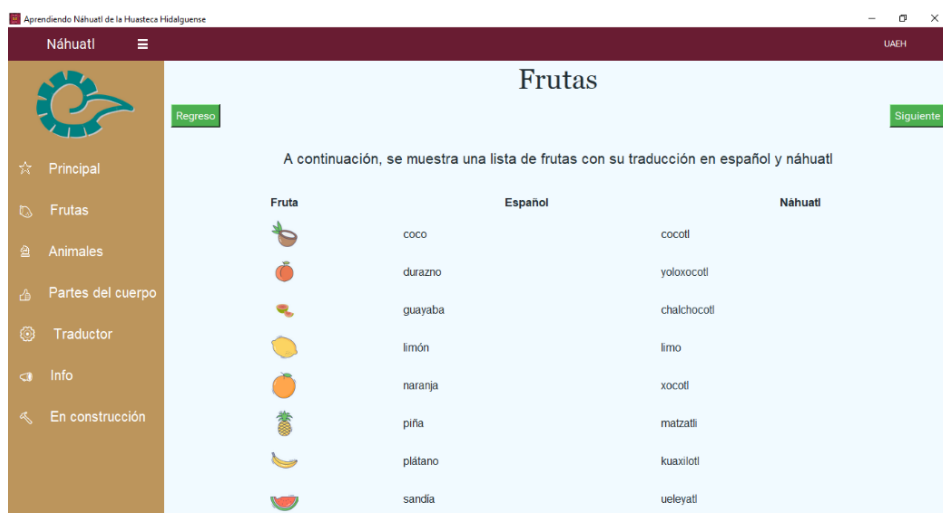


Fig. 6. Contenido del apartado de frutas de la aplicación de escritorio para aprendizaje de Náhuatl.

Puede observarse que en algunos casos, no existe una correspondencia lineal entre cada palabra del Español al Náhuatl, por lo que, es importante tener una clase donde se explica la gramática de la lengua indígena. Finalmente, se presenta en la Fig. 7 la ventana que despliega el botón del juego de memorama en la sección de frutas.

Aquí la aplicación despliega dicha imagen donde es posible memorizar los nombres, colores y frutas. Se considera que la aplicación es hecha para principiantes, es decir, participantes que no saben nada o su conocimiento en la lengua es mínima, por lo que, a las imágenes de las frutas se anexan los nombres en Náhuatl y Español.



Fig. 7. Memorama de frutas de la aplicación para aprendizaje de Náhuatl.

Una vez que los jugadores están listos, pueden dar clic en el botón azul en la parte inferior para iniciar el juego y desplegará la ventana de la Fig. 8. Donde todas las figuras son revueltas de forma aleatoria y pasan a tener el signo de interrogación, que simboliza que no se sabe qué fruta es.



Fig. 8. Memorama en modo incognito de las frutas de la aplicación de escritorio para aprendizaje de Náhuatl.

Cabe mencionarse, que los usuarios pueden empezar a jugar. Una vez que un usuario encuentra dos imágenes iguales, ambas quedan volteadas, en caso contrario, las imágenes vuelven a voltearse. De esta manera hasta terminar el juego, encontrando cada par de imágenes. Al final, automáticamente se vuelve a activar el botón de iniciar el juego, lo que permite inicializarlo las veces que sean necesarias. De manera similar, se desarrolla el apartado de animales (ver Fig. 9) y el de partes del cuerpo (ver Fig. 10). En estos apartados no solo se muestran las palabras correspondientes a cada categoría, sino que también se incluyen ejemplos en forma de oraciones completas, lo cual permite que los participantes no se limiten a memorizar vocabulario aislado.

Al observar cómo estas palabras se integran en estructuras oracionales, los estudiantes pueden comprender con mayor claridad la formulación gramatical y la escritura en lengua Náhuatl, favoreciendo un aprendizaje contextualizado y significativo. Así, el recurso no solo enriquece el conocimiento léxico, sino que también abre la posibilidad de practicar la lengua en situaciones comunicativas reales, fortaleciendo tanto la comprensión como la producción escrita.

Sin embargo, ahora se describe el contenido correspondiente al juego del ahorcado, el cual se plantea como una actividad educativa y entretenida que puede adaptarse para la enseñanza del vocabulario en Náhuatl, en particular en categorías como partes del cuerpo, animales y frutas. Este juego no solo motiva al aprendiente, sino que además se vincula directamente con las palabras utilizadas en el traductor entrenado con la red LSTM, de manera que el vocabulario previamente procesado y vectorizado se convierte también en insumo para una actividad lúdica.

En la preparación, la aplicación de escritorio selecciona de forma aleatoria una palabra secreta relacionada con las categorías mencionadas. Dicha palabra aparece representada por una serie de guiones bajos que corresponden a cada letra. Para iniciar, el jugador intenta adivinar la palabra sugiriendo letras una a una. Si la letra está en la palabra, esta se revela en todas sus posiciones; si no, se dibuja progresivamente una parte del personaje del ahorcado (soporte, cuerda, cabeza, cuerpo, brazo derecho, brazo izquierdo, pierna derecha y pierna izquierda), hasta un máximo de cinco errores permitidos.

Finalmente, el juego termina cuando el jugador logra adivinar la palabra completa o cuando se completa el dibujo del hombre ahorcado, es decir, gana o pierde. En la Fig. 11 se muestra un ejemplo del juego aplicado a la categoría de partes del cuerpo.

Imagen	Oración	Traducción en náhuatl
	El perro juega la pelota	Chichi auitia ica
	La vaca comerá pasto	Kuakuhue tlakuas sakatl
	La gallina es muy bonita	Totolin tlahuēl kuātsin
	La mariposa vuela porque tiene alas	Papalotl patlani pampa ki pia imastlakapal
	El gato está durmiendo	Miston kochtok
	El conejo brinca mucho	Tochtli chitoni miak
	hormiga	ashkatl

Fig. 9. Contenido del apartado de animales de la aplicación de escritorio para aprendizaje de Náhuatl.

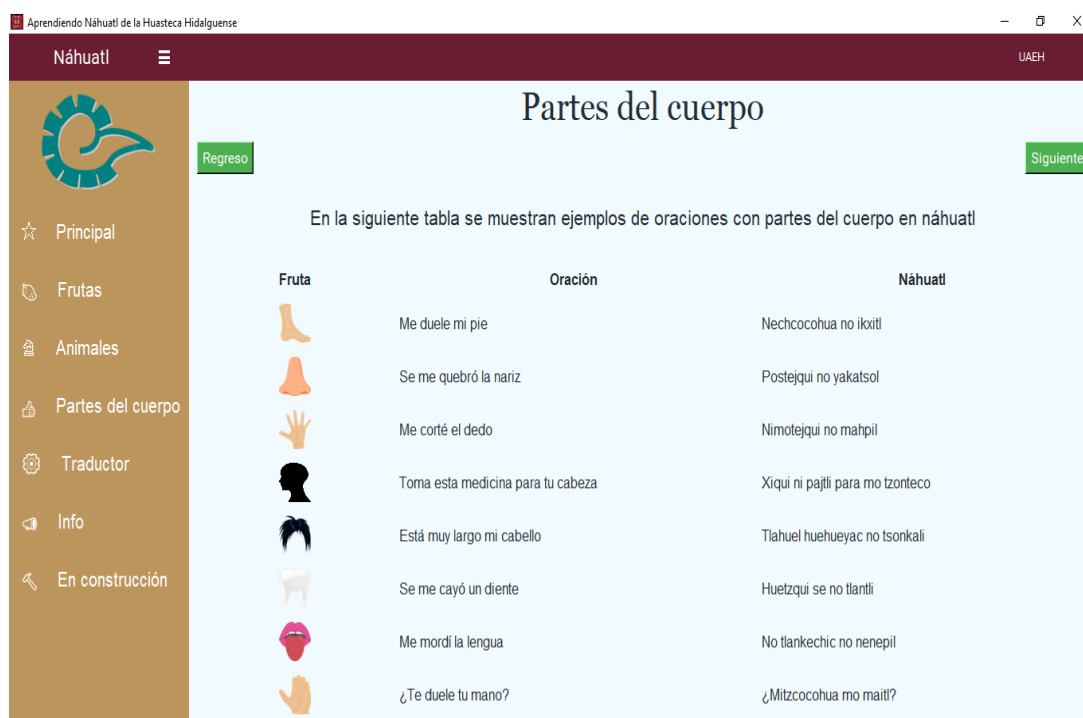


Fig. 10. Contenido del apartado de partes del cuerpo de la aplicación de escritorio para aprendizaje de Náhuatl.

Puede observarse que el usuario ha cometido dos errores, enlistando las letras incorrectas en la parte inferior de los guiones. Gracias a esta integración, el vocabulario que la red neuronal LSTM procesa para la traducción se reutiliza en el juego, reforzando de manera práctica y divertida el aprendizaje del Náhuatl.

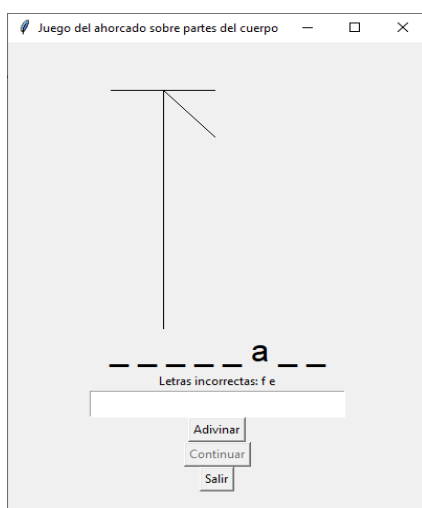


Fig. 11. Juego del ahorcado para la sección de partes del cuerpo de la aplicación de escritorio para aprendizaje de Náhuatl.

Por otro lado, si el usuario adivina la letra será mostrada como se presenta en la Fig. 11 (para la letra a). En caso contrario de que el usuario sepa la palabra la puede escribir en el recuadro blanco y puede darle clic en Adivinar, si es el caso la palabra aparecerá, en caso contrario, el usuario debe seguir jugando y dar clic en Continuar. Finalmente, el usuario puede salir del juego simplemente dando clic a la pestaña de Salir.

Por otro lado, en el apartado del traductor se muestra la Fig. 12. Donde dando clic en el apartado verde de traductor la aplicación de escritorio desplegará automáticamente la ventana del traductor de la Fig. 2.

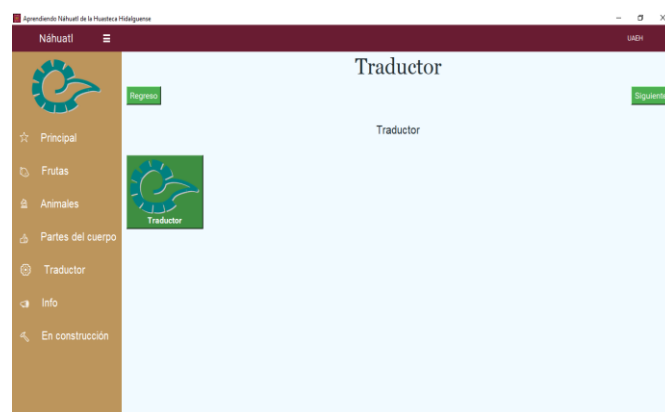


Fig. 12. Sección de la aplicación de escritorio donde se encuentra el traductor de Náhuatl.

En esta ventana el usuario podrá insertar palabras contenidas en el corpus contemplado para la aplicación de escritorio. Es importante mencionar que pueden traducirse palabras del Español al Náhuatl y viceversa, en caso de no contener la palabra a traducir la aplicación dará un mensaje de palabra no encontrada. A continuación, es presentada la discusión de los resultados obtenidos con un grupo de 35 alumnos de 6to grado de la Escuela Primaria Fuentes y Bravo con clave 13DPR0863E, ubicada en la colonia Nueva Francisco I. Madero, Pachuca de Soto, Hidalgo, México.

5. **Discusión de los resultados**

Los resultados obtenidos son resultado de una sesión de uso de la aplicación de escritorio con alumnos de primaria de entre 11 y 13 años. Donde se realizó una evaluación preTest y posTest para conocer el aporte de la aplicación a nivel aprendizaje y un posTest para evaluar la escala de usabilidad, la facilidad de uso, el atractivo visual, el contenido educativo y la motivación y satisfacción de los participantes de la aplicación de escritorio. Para ello, se consideró el siguiente preTest antes de la intervención con la aplicación cuyo objetivo es conocer el conocimiento previo del estudiante sobre palabras en Náhuatl incluyendo las siguientes preguntas:

Instrucciones: Indica tu nivel de acuerdo con las siguientes afirmaciones utilizando la escala de Likert (como un instrumento inicial) considerada como totalmente en desacuerdo 1 (TD), en desacuerdo 2 (ED), neutral 3 (N), de acuerdo 4 (DA) y totalmente de acuerdo 5 (TD).

1. Conozco algunas palabras en Náhuatl relacionadas con partes del cuerpo.
2. Reconozco palabras en Náhuatl que se refieren a animales.
3. Sé cómo se dice 'fruta' en Náhuatl.
4. Puedo identificar palabras en Náhuatl que nombran frutas comunes.
5. Tengo conocimiento de palabras en Náhuatl relacionadas con animales locales.

Tabla 2. Resultados obtenidos del preTest usando escala de Likert.

	TD	ED	N	DA	TA
1. Conozco varias palabras en Náhuatl relacionadas con partes del cuerpo.	27	15	0	0	0
2. Reconozco palabras en Náhuatl que se refieren a animales.	25	17	0	0	0
3. Sé cómo se dice 'fruta' en Náhuatl.	24	18	0	0	0
4. Puedo identificar palabras en Náhuatl que nombran frutas comunes.	29	13	0	0	0
5. Tengo conocimiento de palabras en Náhuatl relacionadas con animales	30	12	0	0	0

Obteniendo los siguientes resultados de la Tabla 2, donde la mayoría de los alumnos respondió que totalmente en desacuerdo y/o en desacuerdo en que conoce, reconoce, sabe, identifica o tiene el conocimiento de palabras en Náhuatl, tal como puede observarse en la Fig. 13, donde prevalece la tendencias de totalmente en desacuerdo (color azul marino) y en desacuerdo (color naranja).

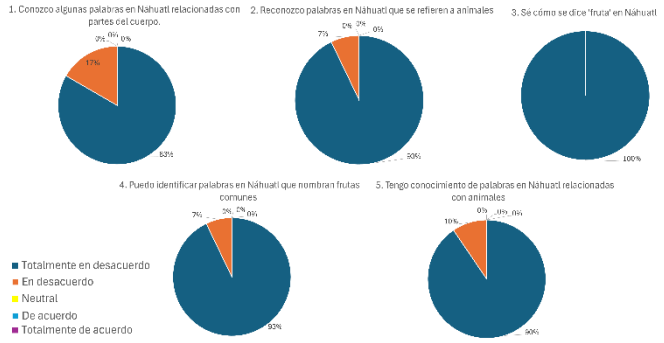


Fig. 13. Resultados de la prueba preTest usando la escala de Likert para los participantes de Náhuatl.

Posteriormente al preTest, se realizó la actividades con los alumnos donde experimentaron con la aplicación de escritorio explorando cada una de las secciones mencionadas anteriormente. Donde pudieron experimentar con los juegos, el traductor y el contenido de la aplicación. Finalizando la actividad se realizó el posTest cuyo objetivo es evaluar el conocimiento adquirido del estudiante sobre palabras en Náhuatl mediante el uso de la aplicación como se presenta a continuación (ver Fig. 14).

Instrucciones: Indica tu nivel de acuerdo con las siguientes afirmaciones utilizando la escala de Likert considerada como totalmente en desacuerdo 1, en desacuerdo 2, neutral 3, de acuerdo 4 y totalmente de acuerdo 5 obteniendo los resultados de la Tabla 3.

1. Ahora conozco varias palabras en Náhuatl relacionadas con partes del cuerpo.
2. Reconozco palabras en Náhuatl que se refieren a animales.
3. Sé cómo se dice 'fruta' en Náhuatl.
4. Puedo identificar palabras en Náhuatl que nombran frutas comunes.
5. Tengo conocimiento de palabras en Náhuatl relacionadas con animales locales.

Tabla 3. Resultados obtenidos del posTest usando escala de Likert.

	TD	ED	N	DA	TA
1. Ahora conozco algunas palabras en Náhuatl relacionadas con partes del cuerpo.	0	0	0	7	35
2. Reconozco palabras en Náhuatl que se refieren a animales.	0	0	0	3	39
3. Sé cómo se dice 'fruta' en Náhuatl.	0	0	0	0	42
4. Puedo identificar palabras en Náhuatl que nombran frutas comunes.	0	0	0	3	39
5. Tengo conocimiento de palabras en Náhuatl relacionadas con animales	0	0	0	4	38

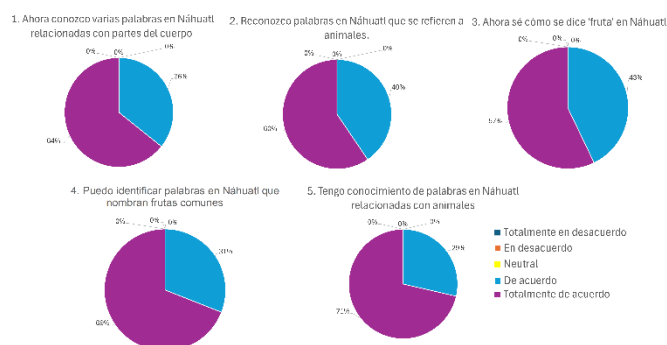


Fig. 14. Resultados de la prueba posTest usando la escala de Likert para los participantes de Náhuatl

En la Fig. 13 se muestran los resultados iniciales del conocimiento de los estudiantes sobre vocabulario en Náhuatl. Puede observarse que la mayoría manifiesta estar “totalmente de acuerdo” en conocer algunas palabras relacionadas con partes del cuerpo (ítem 1), mientras que en categorías como animales (ítem 2) y frutas (ítem 4) también predominan respuestas positivas, aunque con una ligera presencia de respuestas neutrales y en desacuerdo. Esto evidencia que los participantes tienen un conocimiento previo limitado pero existente, principalmente enfocado en palabras comunes o de uso frecuente. En particular, la identificación de la palabra “fruta” en Náhuatl (ítem 3) presenta unanimidad en el reconocimiento, lo cual indica que ciertos vocablos más básicos ya eran familiares antes de la intervención.

En contraste, la Fig. 14 refleja los resultados posteriores a la implementación del recurso didáctico. Se observa un incremento notable en las respuestas positivas (de acuerdo y totalmente de acuerdo) en todos los ítems, destacando un cambio significativo en el conocimiento de palabras relacionadas con animales y frutas (ítems 2 y 4), que anteriormente presentaban mayor dispersión. Asimismo, en los cinco ítems se aprecia que las respuestas tienden a concentrarse en niveles altos de acuerdo, lo cual sugiere que la actividad tuvo un impacto directo en la ampliación del vocabulario en Náhuatl de los estudiantes.

De manera general, la comparación entre ambas figuras muestra que el recurso utilizado favoreció la transición desde un conocimiento parcial o básico hacia un reconocimiento más amplio y seguro de las palabras en Náhuatl en diferentes categorías. Esto confirma la eficacia del material y de las actividades implementadas (como el traductor y el juego del ahorcado) para fortalecer el aprendizaje de la lengua.

Asimismo, fue evaluada la efectividad de la aplicación y la experiencia del usuario mediante una escala de usabilidad del sistema (SUS). La SUS es una herramienta estandarizada ampliamente utilizada para medir la usabilidad percibida de un sistema o aplicación, proporcionando una puntuación que refleja qué tan fácil y satisfactorio resulta para los usuarios interactuar con la herramienta. Esta escala permite identificar fortalezas y áreas de mejora en el diseño y la funcionalidad, contribuyendo a optimizar la experiencia de aprendizaje.

La SUS consiste en 10 afirmaciones que los participantes valoran mediante una escala Likert de 5 puntos, que va desde “totalmente en desacuerdo” hasta “totalmente de acuerdo”. Las afirmaciones incluidas en la evaluación fueron:

1. Creo que me gustaría usar este sistema con frecuencia.
2. Encontré el sistema innecesariamente complejo.
3. Pensé que el sistema era fácil de usar.
4. Creo que necesitaría la ayuda de un experto para usar este sistema.
5. Encontré que las diferentes funciones del sistema estaban bien integradas.
6. Pensé que había demasiada inconsistencia en el sistema.
7. Imagino que la mayoría de las personas aprenderían a usar este sistema rápidamente.
8. Encontré el sistema muy engorroso de usar.
9. Me sentí confiado usando el sistema.
10. Necesitaría aprender muchas cosas antes de poder usar este sistema correctamente.

La aplicación de esta herramienta permitió recopilar datos cuantitativos sobre la percepción de los usuarios, identificando qué aspectos del sistema resultaron intuitivos y cuáles podrían generar dificultades o frustración (ver Tabla 4).

Los resultados obtenidos, mostrados en la Fig. 15, indican tendencias positivas en cuanto a la facilidad de uso y la confianza al interactuar con la aplicación, aunque también señalan oportunidades de mejora en términos de consistencia de funciones y complejidad percibida.

Tabla 4. Resultados obtenidos de usabilidad con escala de Likert.

	TD	ED	N	DA	TA
1. Creo que me gustaría usar este sistema con frecuencia.	0	1	3	20	18
2. Encontré el sistema innecesariamente complejo.	21	15	5	1	0
3. Pensé que el sistema era fácil de usar.	1	1	5	15	20
4. Creo que necesitaría la ayuda de un experto para usar este sistema.	20	18	3	0	1
5. Encontré que las diferentes funciones del sistema estaban bien integradas.	0	0	2	12	28
6. Pensé que había demasiada inconsistencia en el sistema.	18	20	3	0	1
7. Imagino que la mayoría de las personas aprenderían a usar este sistema rápidamente.	0	0	2	12	28
8. Encontré el sistema muy engorroso de usar.	25	15	1	1	0
9. Me sentí confiado usando el sistema.	0	3	4	10	25
10. Necesitaría aprender muchas cosas antes de poder usar este sistema correctamente.	20	15	5	2	0

Este tipo de evaluación es fundamental, ya que la usabilidad influye directamente en la motivación y el aprendizaje de los estudiantes, especialmente en aplicaciones educativas dirigidas a niños. Una herramienta que resulta intuitiva y agradable de usar aumenta la probabilidad de que los estudiantes exploren más contenidos, repitan actividades y desarrollen habilidades lingüísticas de manera autónoma y sostenida. En consecuencia, los hallazgos de la SUS no solo sirven para medir la satisfacción del usuario, sino que también orientan futuras mejoras de diseño que refuercen el valor educativo y cultural de la aplicación.

	TD	ED	N	DA	TA
Facilidad de uso:					
La interfaz de la aplicación es intuitiva	0	0	2	19	21
Pude navegar por la aplicación sin dificultad	0	0	4	18	20
Las instrucciones dentro de la aplicación fueron claras	0	0	3	14	25
Atractivo visual:					
El diseño visual de la aplicación es atractivo	0	0	4	17	21
Los colores y gráficos utilizados en la aplicación son agradables	1	2	5	18	16
Contenido educativo:					
Las actividades de aprendizaje fueron útiles para recordar palabras en náhuatl	0	0	3	22	17
La aplicación ofrece una variedad adecuada de ejercicios	0	2	5	17	18
Motivación y satisfacción:					
Me sentí motivado para continuar usando la aplicación	0	1	3	18	20
Estoy satisfecho con mi experiencia general con la aplicación	0	2	5	18	17
Recomendaría la aplicación a otras personas interesadas en aprender Náhuatl	0	0	2	17	23

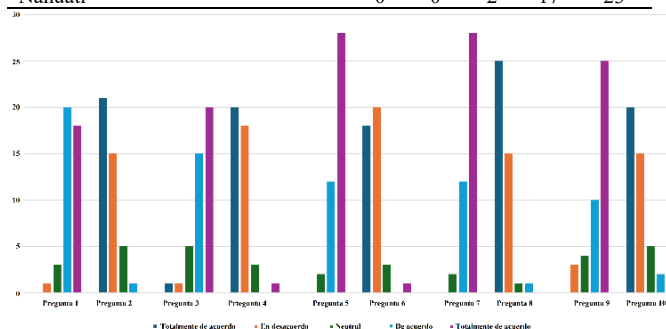


Fig. 15. Resultados de la escala de usabilidad del sistema (SUS) usando la escala de Likert para los participantes de Náhuatl.

Por otro lado, para calcular la puntuación SUS, es necesario restar 1 al valor de las preguntas impares (1, 3, 5, 7, 9) y restar la puntuación de las afirmaciones pares (2, 4, 6, 8, 10) de 5. Posteriormente se suman todos los valores obtenidos y se multiplica el total por 2.5 para obtener una puntuación final entre 0 y 100. Realizando los cálculos correspondientes se tiene un resultado de 84.7. Por lo que, según la escala de SUS los resultados de los usuarios arrojan una buena usabilidad, que se considera que probablemente recomienden el sistema.

Finalmente, se consideran algunas preguntas adicionales que consideran que están diseñadas para evaluar aspectos específicos de la aplicación en cuanto a temas de facilidad de uso, atractivo visual, contenido educativo y motivación y satisfacción.

Facilidad de uso:

La interfaz de la aplicación es intuitiva.
Pude navegar por la aplicación sin dificultad.
Las instrucciones dentro de la aplicación fueron claras.

Atractivo visual:

El diseño visual de la aplicación es atractivo.
Los colores y gráficos utilizados en la aplicación son agradables.

Contenido educativo:

Las actividades de aprendizaje fueron útiles para recordar palabras en Náhuatl."
La aplicación ofrece una variedad adecuada de ejercicios.

Motivación y satisfacción:

Me sentí motivado para continuar utilizando la aplicación.
Estoy satisfecho con mi experiencia general con la aplicación.
Recomendaría esta aplicación a otras personas interesadas en aprender Náhuatl.

Los datos obtenidos y sintetizados en la Tabla 5 evidencian una valoración positiva de la aplicación en los cuatro ejes principales: facilidad de uso, atractivo visual, contenido educativo y motivación/satisfacción.

Tabla 5. Resultados de preguntas adicionales usando escala de Likert.

Por otro lado, la Fig. 16 presenta los resultados gráficos para facilidad de uso la mayoría de los estudiantes respondieron que están de acuerdo y totalmente de acuerdo con la interfaz, la navegación y las instrucciones e indicaciones que tiene la misma.

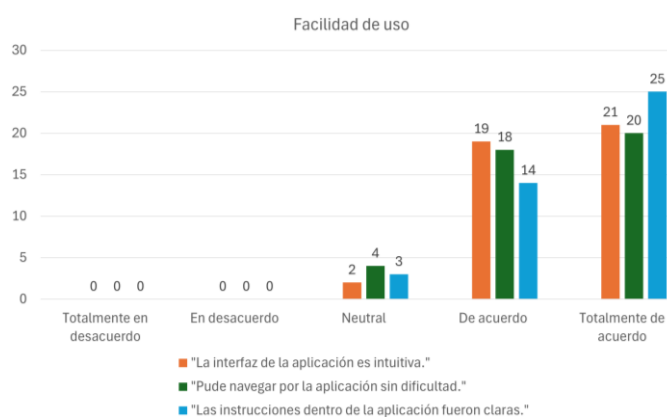


Fig. 16. Resultados de la prueba de facilidad de uso usando la escala de Likert para los participantes de Náhuatl.

En cuanto al atractivo visual de la aplicación al menos 3 usuarios respondieron que el diseño visual de la aplicación no lo consideran atractivo y que los colores y gráficos utilizados en la aplicación son poco agradables. En el caso particular de este punto, preguntando directamente a los usuarios consideraron que usar colores como el guinda no ayudan mucho en la aplicación de escritorio, por lo que, se buscará a futuro mejorar e incluso cambiarla por colores más claras. A pesar del detalle de los colores y el diseño visual los resultados muestran que a la mayoría de los usuarios le gusto

y no consideran que los colores sean algún inconveniente en el diseño y uso de esta. Finalmente, respecto al diseño visual los resultados reflejan que al menos a 28 participantes les agrado o están de acuerdo y complementamente de acuerdo con el diseño (ver Fig. 17).

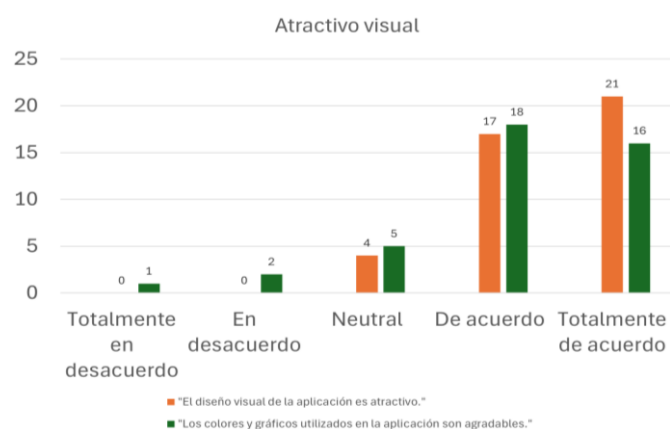


Fig. 17. Resultados de la prueba de atractivo visual usando la escala de Likert para los participantes de Náhuatl.

En cuanto al contenido educativo, los resultados resaltan que al menos a 35 personas dichas actividades de aprendizaje fueron útiles para recordar y aprender algunas palabras en Náhuatl. Por otro lado, también puede mencionarse que con base en los usuarios la aplicación ofrece una variedad adecuada de ejercicios y ejemplos que ayudan a aprender la lengua. La Fig. 18 muestra que la mayoría de los estudiantes seleccionaron estar de acuerdo y totalmente de acuerdo en el contenido educativo.

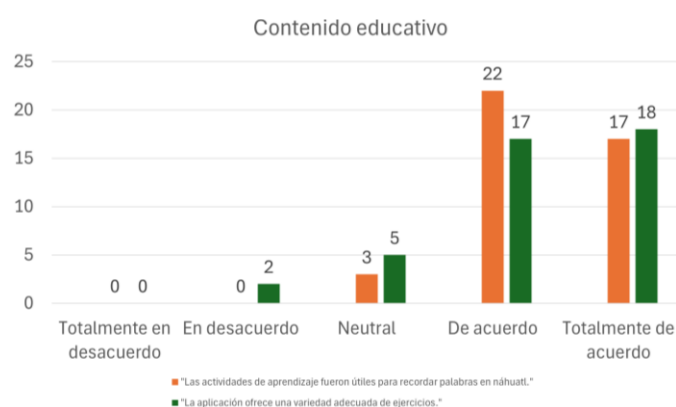


Fig. 18. Resultados de la prueba de contenido educativo usando la escala de Likert para los participantes de Náhuatl.

Finalmente, respecto a la motivación y la satisfacción de la aplicación de escrito predominan las respuestas de acuerdo y totalmente de acuerdo. Es decir, los resultados reflejan que los usuarios se sintieron motivados para continuar utilizando la aplicación (ver Fig. 19). Incluso, al final de las experiencia muchos mencionaron estar satisfecho con mi experiencia general con la aplicación reflejado en la segunda pregunta. Para finalizar, se preguntó a los usuarios si recomendarían la aplicación a otras personas interesadas en aprender Náhuatl,

obteniendo resultados prevalecientes para de acuerdo y totalmente de acuerdo.

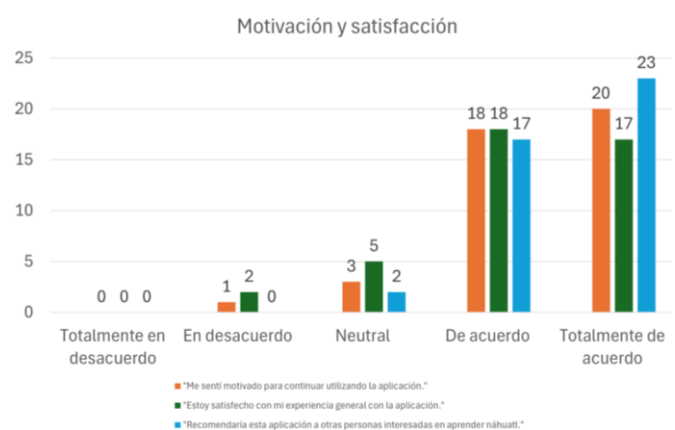


Fig. 19. Resultados de la prueba de motivación y satisfacción usando la escala de Likert para los participantes de Náhuatl.

Por otro lado, se realizó una prueba de usabilidad cualitativa mediante entrevistas semiestructuradas y observación directa con el fin de conocer la experiencia de uso y las percepciones de estudiantes de primaria de 11 a 13 años. Este tipo de evaluación cualitativa se centra en comprender en profundidad la interacción de los usuarios con la aplicación, más allá de métricas numéricas, priorizando la interpretación de sus comentarios, emociones y conductas durante la sesión (Creswell & Poth, 2018).

La metodología aplicada combinó dos técnicas principales. Primero, entrevistas breves semiestructuradas: realizadas al finalizar la interacción, permitieron explorar de manera flexible las opiniones de los estudiantes sobre la facilidad de uso, la claridad de las instrucciones y la motivación generada por la aplicación. Segundo, observación directa no participante: durante el uso de la herramienta, se registraron expresiones verbales y no verbales, así como las dificultades que surgían al interactuar con la interfaz. Este enfoque posibilitó captar información que los estudiantes no siempre expresan de manera explícita, pero que se refleja en sus gestos, pausas o dudas (Nielsen, 1994).

La combinación de entrevistas y observación facilitó recopilar información rica y contextualizada, lo que permitió identificar tanto los aspectos más atractivos de la aplicación (como el valor lúdico y la motivación) como las posibles áreas de mejora (por ejemplo, la necesidad de simplificar instrucciones o reforzar ayudas visuales). Este método cualitativo resulta especialmente apropiado en contextos educativos con niños, ya que prioriza la comprensión de la experiencia subjetiva del usuario en lugar de limitarse a indicadores técnicos de desempeño (Rubin & Chisnell, 2008).

Algunos hallazgos positivos revelaron que la aplicación logra generar motivación entre los participantes. Donde actividades como el juego del ahorcado resultaron entretenidas y fomentaron el interés por seguir aprendiendo Náhuatl. Asimismo, los estudiantes valoraron la utilidad del traductor integrado, ya que les permitió comprobar palabras y mejorar la ortografía, reforzando su aprendizaje de manera práctica. Otro aspecto destacado fue la retroalimentación inmediata: los avisos de errores, tanto en letras como en traducciones, ayudaron a los niños a corregirse en el momento, lo que contribuyó a un aprendizaje más dinámico y efectivo.

Sin embargo, también se identificaron problemas que afectan la experiencia de uso. Algunos estudiantes manifestaron dificultad para localizar ciertas secciones de la aplicación, como el vocabulario de frutas, lo que generó la necesidad de solicitar apoyo externo. Además, se observó que cuando el traductor no reconocía una palabra, los niños requerían ejemplos o sugerencias adicionales para poder avanzar en sus ejercicios. En cuanto al juego del ahorcado, varios participantes expresaron frustración al no poder adivinar palabras desconocidas, prolongándose los tiempos de resolución y aumentando los errores.

Los comentarios representativos reflejan estas experiencias: un estudiante comentó, *“Me gustó porque adivinar palabras es como jugar con mis amigos”*, mientras que otro señaló, *“A veces no sabía dónde buscar la lista de frutas, tuve que pedir ayuda”*. Estos testimonios resaltan tanto los aspectos motivadores como las barreras que encontraron los usuarios durante la prueba.

Los resultados cualitativos sugieren que la aplicación tiene un alto potencial educativo al favorecer la motivación y el aprendizaje del vocabulario en Náhuatl. Los niños reconocen haber aprendido nuevas palabras de manera divertida, lo que indica que el componente lúdico es efectivo. No obstante, los hallazgos también evidencian áreas de mejora importantes, especialmente en la navegación interna, la provisión de ayudas contextuales en el traductor y la adaptación del nivel de dificultad del juego del ahorcado. Estas observaciones apuntan a la necesidad de realizar ajustes que garanticen una experiencia de uso más accesible y menos frustrante, asegurando que los beneficios educativos se maximicen para todos los usuarios.

Por lo tanto, con los resultados obtenidos se puede demostrar que la aplicación de escritorio tiene un gran potencial para emplearse como una herramienta educativa efectiva. Su diseño y funcionalidad permiten facilitar el aprendizaje de manera interactiva y atractiva para los usuarios. Sin embargo, para maximizar su impacto, es importante considerar algunos aspectos que fueron sugeridos en los comentarios de los usuarios.

En primer lugar, la incorporación de detalles más cuidadosos en la selección y combinación de colores podría mejorar significativamente la experiencia visual, haciendo que el entorno de aprendizaje sea más agradable y estimulante. Además, es recomendable enriquecer el contenido con una mayor variedad de actividades que promuevan diferentes formas de aprendizaje, como ejercicios prácticos, cuestionarios interactivos o retos creativos.

Asimismo, durante la prueba de la aplicación de escritorio (ver Fig. 20) varios estudiantes mencionaron la posibilidad de que se puedan integrar más elementos multimedia que hagan la experiencia aún más dinámica y entretenida. Por ejemplo, la inclusión de animaciones podría ayudar a ilustrar conceptos complejos de manera más clara, mientras que la adición de juegos educativos serviría para reforzar los conocimientos adquiridos a través de la diversión. También se mencionó la utilidad de incorporar audios, ya sea para explicar instrucciones, narrar textos o añadir efectos sonoros y más juegos que permitan aprender.

En conjunto, estas mejoras no solo podrían enriquecer la aplicación, sino también aumentar el grado de motivación y compromiso de los estudiantes, convirtiéndola en una herramienta más completa y adaptada a las necesidades

educativas actuales. Por lo tanto, se recomienda continuar desarrollando estas funcionalidades y realizar pruebas con grupos más amplios para validar su efectividad en diferentes contextos de aprendizaje.



Fig. 20. Pruebas de la aplicación de escritorio con alumnos de 6to grado de la Escuela Primaria Fuentes y Bravo.

6. Conclusiones

La aplicación de escritorio para el aprendizaje del Náhuatl demuestra un alto potencial educativo y cultural. Los resultados cualitativos evidencian que los estudiantes de primaria encuentran motivadora la experiencia, especialmente mediante juegos interactivos como el ahorcado y el memorama, que favorecen la participación activa y la retención del vocabulario. La retroalimentación inmediata y el traductor integrado contribuyen a un aprendizaje más autónomo y efectivo, reforzando la ortografía y la comprensión de las palabras en Náhuatl. La evaluación cuantitativa, mediante cuestionarios con escalas de acuerdo, confirma que la mayoría de los usuarios logró reconocer y recordar palabras básicas relacionadas con frutas y verduras, animales, partes del cuerpo y oraciones con estas clases de palabras, validando la efectividad de la combinación de juegos, actividades interactivas y el traductor.

Además de su valor pedagógico, la aplicación constituye un aporte significativo a la preservación y promoción del idioma Náhuatl, centrando su corpus en palabras básicas que permiten un acercamiento práctico y contextualizado, apoyando la revitalización lingüística y fomentando el interés de las nuevas generaciones por su patrimonio cultural. La integración de redes neuronales LSTM en el traductor representa un avance innovador en inteligencia artificial educativa, mejorando la precisión de las traducciones y demostrando cómo la tecnología puede potenciar la enseñanza de lenguas originarias.

Los estudiantes también muestran una mayor motivación para explorar otros aspectos culturales asociados al idioma, como canciones, refranes y tradiciones locales, lo que evidencia que la aplicación no solo enseña vocabulario, sino que también fortalece la identidad cultural y el sentido de pertenencia. Esta experiencia integral fomenta la curiosidad, la interacción y la apreciación del patrimonio lingüístico de manera lúdica y significativa. En conjunto, los hallazgos cualitativos y cuantitativos muestran que la aplicación no solo facilita el

aprendizaje del Náhuatl de manera divertida y accesible, sino que también promueve la valoración y el uso activo de la lengua, evidenciando que la innovación tecnológica puede ser un aliado poderoso en la preservación de idiomas indígenas.

Finalmente, como línea de trabajo futuro se plantea fortalecer la aplicación de escritorio mediante la incorporación de nuevos juegos y dinámicas que integren de forma directa el traductor. Se propone además combinar métodos cuantitativos y cualitativos, implementar evaluaciones de rendimiento técnico y realizar pruebas de usuario con métricas objetivas. Esto incluirá el análisis del tiempo de ejecución de las tareas, la tasa de errores y otros indicadores que permitan obtener una retroalimentación precisa y enriquecer el diseño, garantizando una experiencia de aprendizaje más completa y efectiva.

Agradecimientos

Los autores de este artículo expresan su más sincero agradecimiento a la Directora Ángela Monroy Hernández y al profesor Edmundo Gómez de la Escuela Primaria Fuentes y Bravo, con clave 13DPR0863E, ubicada en la colonia Nueva Francisco I. Madero en Pachuca de Soto, por las facilidades otorgadas y el valioso apoyo brindado para la realización de las pruebas piloto con los alumnos del sexto grado, grupo “A”, durante el ciclo escolar 2024-2025.

Su colaboración fue fundamental para llevar a cabo la validación de la aplicación educativa presentada en este trabajo, la cual tiene como propósito fomentar el aprendizaje y la valoración de las lenguas indígenas, en particular del náhuatl. Gracias a la disposición y entusiasmo tanto del personal docente como de los estudiantes, fue posible generar un entorno de aprendizaje dinámico y participativo que enriqueció el proceso de evaluación.

Referencias

- Afifah, N., Yulita, S., & Sarathan, M. (2021). Natural Language Processing for Automatic Sentiment Analysis in Social Media Data. *International Journal of Information Engineering and Science*, 1(1), 16–19. <https://doi.org/10.62951/ijies.v1i2.54>
- Bautista Morales, R., Martínez Ramírez, Y., Rocha Peña, L. E., & Montes Santiago, R. E. (2024). Arquitectura de un traductor automático para el idioma mixteco: un enfoque específico para lenguas indígenas con escasos recursos lingüísticos. *Revista De Investigación En Tecnologías De La Información*, 12(28 Especial), 71–81. <https://doi.org/10.36825/RITI.12.28.007>
- Brownlee, J. (2020). Deep learning for time series forecasting: predict the future with MLPs, CNNs and LSTMs in Python. *Machine Learning Mastery*.
- Çelik, T., & Koç, E. (2021). Natural Language Processing Techniques for Text Classification of Biomedical Documents: A Systematic Review. *Information*, 13(10), 499. <https://doi.org/10.3390/info13100499>
- Creswell, J. W., & Poth, C. N. (2018). *Qualitative Inquiry & Research Design: Choosing Among Five Approaches* (4th ed.). SAGE Publications.
- Devlin, J., Chang, M.-W., Lee, K., & Toutanova, K. (2019). BERT: Pre-training deep bidirectional transformers for language understanding. *arXiv*. <https://doi.org/10.48550/arXiv.1810.04805>
- Gers, F. A., Schmidhuber, J., & Cummins, F. (2000). Learning to forget: Continual prediction with LSTM. *Neural computation*, 12(10), 2451–2471. <https://doi.org/10.1162/089976600300015015>
- Gers, F. A., Schraudolph, N. N., & Schmidhuber, J. (2002). Learning precise timing with LSTM recurrent networks. *Journal of machine learning research*, 3(Aug), 115–143. <https://doi.org/https://doi.org/10.1162/153244303768966139>
- González, Servín C. (2022). Traductor automático de la lengua purépecha al español usando redes transformer, Tesis de Licenciatura, Universidad Michoacana de San Nicolás de Hidalgo. http://bibliotecavirtual.dgb.umich.mx:8083/xmlui/handle/DGB_UMICH/12218
- Gutiérrez-Vasques, X., Sierra, G., & Hernández Pompa, I. (2016). Axolotl: A web accessible parallel corpus for Spanish-Nahuatl. *Proceedings of the Tenth International Conference on Language Resources and Evaluation (LREC 2016)*, 4210–4214.
- HaCohen-Kerner, Y., Miller, D., & Yigal, Y. (2020). The influence of preprocessing on text classification using a bag-of-words representation. *PLOS ONE*, 15(5), e0232525. <https://doi.org/10.1371/journal.pone.0232525>
- Herrera-Ortiz, J., Peña-Avilés, J., Herrera-Valdivieso, M., Moreno-Morán, D. (2024). La inteligencia artificial y su impacto en la comunicación: recorrido y perspectivas. *Telos: Revista de Estudios Interdisciplinarios en Ciencias Sociales*, 26(1), 278–296. www.doi.org/10.36390/telos261.18
- Hidalgo-Ternero, R. (2020). Traducción automática neuronal: Un enfoque para mejorar la calidad y fluidez en la traducción de lenguas indígenas. *Revista de Lingüística y Lenguas Aplicadas*, 15(2), 45–58.
- INEGI. (2024). Estadísticas a propósito del día internacional de los pueblos indígenas, Comunicado de prensa núm. 472/24. Disponible: mayo 2025. https://www.inegi.org.mx/contenidos/saladeprensa/aproposito/2024/EAP_PueblosInd24.pdf
- Javed, A., Zan, H., Mamrybayev, O., Abdullah, M., Ahmed, K., Oralbekova, D., Dinara, K. y Akhmediyarova, A. (2025). Modelo de reordenamiento basado en transformadores para mejorar la traducción contextual y sintáctica en la traducción automática neuronal de bajos recursos. *Electronics*, 14 (2), 243. <https://doi.org/10.3390/electronics14020243>
- Jia, Y., Wu, Z., Xu, Y., Ke, D., & Su, K. (2017). Long short-term memory projection recurrent neural network architectures for piano’s continuous note recognition. *Journal of Robotics*, 2017, 1–7. <https://doi.org/https://doi.org/10.1155/2017/2061827>
- Kalyan, KS; Rajasekharan, A.; Sangeetha, S. Ammus. (2021). Un estudio de modelos preentrenados basados en transformadores en el procesamiento del lenguaje natural. *ArXiv*:2108.05542.
- Kang, H., Yang, S., Huang, J., & Oh, J. (2020). Time series prediction of wastewater flow rate by bidirectional LSTM deep learning. *International Journal of Control, Automation Systems*, 18, 3023–3030. <https://doi.org/https://doi.org/10.1007/s12555-019-0984-6>
- Kang, L., Di, L., Deng, M., Yu, E., & Xu, Y. (2016). Forecasting vegetation index based on vegetation-meteorological factor interactions with artificial neural network. 2016 Fifth International Conference on Agro-Geoinformatics (Agro-Geoinformatics)
- Kurniawan, R., & Maharani, D. (2020). Natural Language Processing for Automatic Sentiment Analysis in Social Media Data. *International Journal of Information Engineering and Science*, 1(1), 16–19. <https://doi.org/10.62951/ijies.v1i2.54>
- Liang, L., Liu, X., & Zhang, Y. (2020). Challenges in word vectorization for indigenous languages: A case study on Nahuatl. *Proceedings of the 12th International Conference on Language Resources and Evaluation (LREC 2020)*, 1234–1241.
- Nahapetyan, Y. (2019). The benefits of the Velvet Revolution in Armenia: Estimation of the short-term economic gains using deep neural networks. *Central European Economic Journal*, 6(53), 286–303. <https://doi.org/https://doi.org/10.2478/ceej-2019-0018>
- Nielsen, J. (1994). *Usability Engineering*. Morgan Kaufmann Publishers.
- Nurdin, N., & Maharani, D. (2020). Natural Language Processing for Automatic Sentiment Analysis in Social Media Data. *International Journal of Information Engineering and Science*, 1(1), 16–19. <https://doi.org/10.62951/ijies.v1i2.54>
- Reddy, D. S., & Prasad, P. R. C. (2018). Prediction of vegetation dynamics using NDVI time series data and LSTM. *Modeling Earth Systems Environment*, 4, 409–419. <https://doi.org/https://doi.org/10.1007/s40808-018-0431-3>
- Rubin, J., & Chisnell, D. (2008). *Handbook of Usability Testing: How to Plan, Design, and Conduct Effective Tests* (2nd ed.). Wiley Publishing.
- Santiago-Benito, H.; Córdova-Esparza, D.-M.; Castro Sánchez, N.-A.; García-Ramírez, T.; Romero-González, J.-A.; Terven, J. (2024). Automatic Translation from Mixtec to Spanish Languages Using Neural Networks. *Appl. Sci.* 2024, 14, 2958. <https://doi.org/10.3390/app14072958>
- Torres-Moreno, J.-M., Guzmán-Landa, J.-J., Ranger, G., Avendaño Garrido, M. L., Figueroa-Saavedra, M., Quintana-Torres, L., González-Gallardo, C.-E., Linhares Pontes, E., Velázquez Morales, P., & Moreno Jiménez, L.-G. (2024). *\$\pi\$-yalli: Un nouveau corpus pour le nahuatl*. *arXiv*. <https://doi.org/10.48550/arXiv.2412.15821>
- Zacarias, Márquez D., & Meza, Ruiz, I. V. (2021). Traductor automático neuronal ayuuk-español, *Research in Computer Science* 150 (5).