

Species Detection Model Overview

Prepared by the WildSync Tech team. This document describes work at an early prototype stage. It contains no performance claims, deployment counts or partner names that we cannot verify.

Purpose

This overview explains how our species detection work is structured, what it can currently do, and what we are not claiming.

Approach

We treat detection as an audio classification problem over short windows of recorded sound. Candidate windows are extracted from continuous recordings and scored, and only scored segments are kept for human review rather than the full raw archive.

- Continuous field recording, segmented into short analysis windows.
- Model scores each window for target call types.
- Human review remains in the loop for anything used in a conclusion.
- Uncertain segments are surfaced, not silently discarded.

Training data

The model is trained on publicly available annotated call libraries plus our own field recordings. Our own recordings are limited in volume and geography, which is the main constraint on how far the model can currently generalise.

What we do not claim

We do not publish an accuracy figure. A single headline accuracy number would be misleading given the size and narrowness of our evaluation set. When we report performance, it will be per species, with the evaluation set described.

- No claimed species count beyond what we have evaluated.
- No claim of real-time detection on-device.
- No claim of performance in habitats we have not recorded in.

Human oversight

A detection is a prompt for a person to listen, not a conclusion. Any output that informs a management decision is reviewed by a human first.

Questions or corrections: kavanfaulkner@wildsyncai.org