

Ako sa littlEye AI učí — a prečo občas odpovie zle

Tento krátky návod vysvetľuje, ako čítať výsledky modelu, prečo model niekedy seabavedomo povie nezmysel, a ako ho vlastnými snímkami krok za krokom učiť, aby sa zlepšoval presne v tom, čo potrebujete.

1 · Prečo model zhodnotil normálne pľúca ako „Srdce 100 %“

Skúsenosť z praxe: snímka normálnych pľúc z webu → model povedal Srdce 100 %; snímka sivej hepatizácie pri lobárnej pneumónii → Pečeň 81,8 %. To nie je náhodná porucha — sú za tým tri pochopiteľné príčiny:

- Mierka. Ak je model trénovaný na výrezoch v mierke ~20× objektívu (zorné pole ~0,5 mm). Prehľadová snímka celého laloka pri malom zväčšení je pre model úplne neznámy svet — alveoly na nej vyzerajú ako drobná textúra, nie ako štruktúry, ktoré pozná.
- Model MUSÍ vybrať jednu zo svojich tried. „100 %“ neznamená istotu patológa — znamená len „zo všetkého, čo poznám, sa mi to najviac podobá na srdce“. Model nemá možnosť povedať „neviem“ ani „toto som nikdy nevidel“.
- Model pozná len to, čo bolo v tréningu. Ak dataset obsahoval iba normálnu histológiu, Sivá hepatizácia — bezvzdušný lalok hustý bunkami — sa modelu logicky podobá na pečeň: hustá, celulárna, bez alveolárnych priestorov. Model nikdy pneumóniu nevidel, tak vybral najbližšie, čo pozná.

Zlaté pravidlo

Model vie hodnotiť len to, na čom bol trénovaný: rovnaká mierka (~20× objektív; v praxi mikrofotografia pri celkovom zväčšení 100×–200×), rovnaké farbenie (HE) a kvalita mikrofotografie, nie webový JPEG. Čím viac sa vaša snímka líši od tréningových, tým menej výsledok znamená.

2 · Ako učenie funguje — dataset, nie čarovanie

Core ML model sa nedá „dotréňovať“ — hotový model je zapečený. Učenie v littlEye AI preto funguje cez dataset: všetko, čo chcete model naučiť, uložíte ako označenú snímku, a model pretrénujete nanovo z celého datasetu. Trvá to minúty a má to veľkú výhodu: každý nový model má poctivé, merateľné metriky na nezávislom teste.

3 · Učiaca slučka — krok za krokom

1. Vyhodnoťte snímku (Vyhodnotenie snímky → Z fotiek / Zo súborov).
2. Nesúhlasíte s výsledkom? Pod výsledkom je nová sekcia „Naučiť z tejto snímky“: vyberte správny orgán a diagnózu a ťuknite Uložiť do datasetu. Snímka sa uloží aj so správnou odpoveďou.
3. V datasete orgánu snímke v detaile doplňte zväčšenie a farbenie; ak je na nej viac nálezov, vyznačte miesto nálezu výrezom — trénuje sa potom výrez.
4. Keď takto nazbierate aspoň 5 snímok na triedu (odporúčame 30+), otvorte Tréning a model pretrénujte. Starší model s rovnakým názvom môžete zmazať.
5. Skontrolujte test % a konfúznú maticu — uvidíte presne, čo si model ešte pletie, a teda ktoré snímky treba nabudúce pridať.

Toto JE to učenie

Slučka vyhodnot' → oprav → ulož → pretrénuj → over je presne to, ako sa aplikácia učí a pamätá si vaše opravy. Každou iteráciou je model o kúsok viac váš — a cez zdieľanie datasetu (.le3set) sa vaše opravy dostanú aj ku kolegom.

4 · Úpravy diagnóz po natrénovaní

- Pridať novú diagnózu: v sekcii orgánu ťuknite Pridať pri „Diagnózy (triedy na tréning)“ — nová trieda je hneď k dispozícii na označovanie aj učiacu slučku.
- Preradiť snímku: v detaile snímky jednoducho vyberte inú diagnózu — zmena sa prejaví pri najbližšom tréningu.
- Diagnózy sa viažu na model až tréningom. Hotový model má triedy zabezčené — premenovanie či pridanie diagnózy v datasete zmení až nasledujúci natrénovaný model.

5 · Čo reálne čakať od koľkých snímok

Snímok na triedu	Čo čakať
5–10	model sa rozbehne, ale výsledky berte len orientačne
30+	použiteľný pracovný model na vlastných dátach
100+ z rôznych prípadov	stabilnejšie správanie aj na cudzích preparátoch
+ vždy	trieda „Norma“, jedno zväčšenie, jedno farbenie, validácia na nezávislých snímkach

Kde začať bez štartovacieho modelu

Skutočnú diagnostickú silu má len model trénovaný na vašich preparátoch cez učiacu slučku vyššie. Odporúčaný postup: vyberte si jeden orgán, ktorý reálne hodnotíte, naplňte ho vlastnými snímkami (Norma + patológie z vašich pitiev) a natrénujte. Prvý použiteľný model vznikne rýchlejšie, než čakáte — na triedu stačia tri rôzne prípady.

Michiganské modely — zmazané

Modely odvodené z michiganskej zbierky sme zmazali. Požiadali sme University of Michigan o stanovisko a odpoveď znela, že naše použitie je komerčné a licencia CC BY-NC-SA ho nepokrýva — a to aj vtedy, keď snímky nie sú súčasťou aplikácie. Stanovisko sme prijali. To isté sa týka modelov z TCGA-UT. Postup opísaný vyššie funguje rovnako aj bez nich.

MUDr. Ján Šikuta, PhD. · forensika.eu@icloud.com · littlEye AI 2.2.3 · august 2026