

Modele End-to-End pentru Recunoașterea Automată a Vorbirii (ASR)

5.1. Introducere: O Schimbare de Paradigmă

Recunoașterea automată a vorbirii (ASR) a fost, timp de decenii, dominată de o arhitectură hibridă, modulară, compusă din componente distincte: un model acustic, un dicționar de pronunție și un model lingvistic.¹ Această abordare, deși a atins performanțe remarcabile, a implicat procese de antrenare separate și complexe, dependente între module și o necesitate stringentă de expertiză lingvistică și date adnotate la nivel fonetic.

În ultimii ani, o nouă paradigmă a redefinit fundamental domeniul: modelele end-to-end (E2E). Acestea înlocuiesc sistemul hibrid tradițional cu o singură rețea neuronală unificată, antrenată să mapeze direct secvența de intrare (trăsături acustice sau eșantioane audio brute) la secvența de ieșire (caractere, foneme sau cuvinte). Această schimbare a fost catalizată de progresele în învățarea profundă, de creșterea exponențială a puterii de calcul și de disponibilitatea unor volume masive de date audio.

Acest capitol explorează arhitecturile fundamentale care stau la baza sistemelor ASR end-to-end. Vom analiza în detaliu două abordări cheie: modelele bazate pe Connectionist Temporal Classification (CTC) și cele de tip encoder-decoder cu mecanism de atenție. Vom discuta, de asemenea, despre arhitecturile hibride care combină avantajele ambelor abordări și vom introduce conceptul revoluționar de învățare auto-supervizată, exemplificat prin modelul Wav2Vec 2.0. Obiectivul este de a oferi o înțelegere aprofundată a principiilor, avantajelor și limitărilor care definesc stadiul actual al tehnologiei ASR.

5.2. De la Arhitectura Modulară la cea Unificată

Pentru a aprecia impactul modelelor end-to-end, este esențial să revedem structura sistemelor tradiționale. Un sistem ASR hibrid clasic implică un pipeline complex:

1. **Modelul Acustic (AM):** De obicei un HMM-GMM sau, mai recent, un hibrid DNN-HMM, acest modul învață o distribuție probabilistică între secvențele de cadre acustice (ex: MFCC-uri) și unitățile sub-lexicale, precum fonemele.
2. **Dicționarul de Pronunție (Lexicon):** O componentă creată manual sau semi-automat care mapează fiecare cuvânt din vocabular la una sau mai multe secvențe de foneme.
3. **Modelul Lingvistic (LM):** Un model statistic (inițial bazat pe n-gramme, ulterior pe rețele neuronale) care evaluează probabilitatea unei secvențe de cuvinte, impunând constrângeri gramaticale și semantice.

Fiecare dintre aceste componente este antrenată și optimizată separat, necesitând seturi de date etichetate și aliniate cu mare atenție, un proces costisitor și predispus la erori. Antrenarea separată poate duce la propagarea erorilor și la sub-optimalitate globală, deoarece fiecare modul este optimizat pentru un obiectiv local, nu pentru sarcina finală de transcriere.

Modelele end-to-end elimină această separare artificială. Ele învață o mapare directă de la acustică la text, integrând implicit cunoștințele de pronunție și de limbaj într-o singură rețea neuronală. Această abordare oferă avantaje semnificative :

- **Simplitate:** Procesul de antrenare este unificat, necesitând doar perechi de date audio și transcrierile corespunzătoare, fără a fi nevoie de aliniamente fonetice.
- **Performanță:** Prin optimizarea globală a întregului sistem pentru funcția de cost finală (ex: minimizarea erorii de transcriere), modelele E2E pot atinge performanțe superioare.
- **Flexibilitate:** Adăugarea de noi limbi sau dialecte este mult simplificată, deoarece nu mai este necesară crearea manuală a unui dicționar de pronunție.

5.3. Connectionist Temporal Classification (CTC)

Una dintre provocările fundamentale în ASR este problema alinierii: lungimea secvenței de trăsături acustice (T) este, în general, mult mai mare decât lungimea secvenței de caractere din transcriere (L), și nu există o corespondență directă, unu-la-unu, între cadrele audio și caractere. Funcția de pierdere Connectionist Temporal Classification (CTC), introdusă de Graves et al. (2006), a fost o inovație crucială care a rezolvat această problemă pentru modelele end-to-end.

Spre deosebire de abordările tradiționale, care presupun segmentarea manuală sau automată a datelor, CTC maximizează probabilitatea marginală a tuturor aliniamentelor posibile între secvența de intrare (cadre acustice) și secvența de ieșire (transcriere), introducând un simbol special „blank” pentru a gestiona discrepanțele de lungime și pentru a permite repetarea sau omisiunea etichetelor.

Într-un model CTC, rețeaua (de obicei un encoder RNN, LSTM sau Transformer) produce o distribuție de probabilitate pentru fiecare cadru acustic, acoperind toate simbolurile posibile (de exemplu, litere, foneme) plus blank. Decodarea se realizează prin selectarea secvenței cu probabilitate maximă, eliminând simbolurile blank și agregând etichetele consecutive. Algoritmul de antrenare utilizează forward-backward dynamic programming, similar cu HMM, pentru a calcula eficient gradientul pierderii

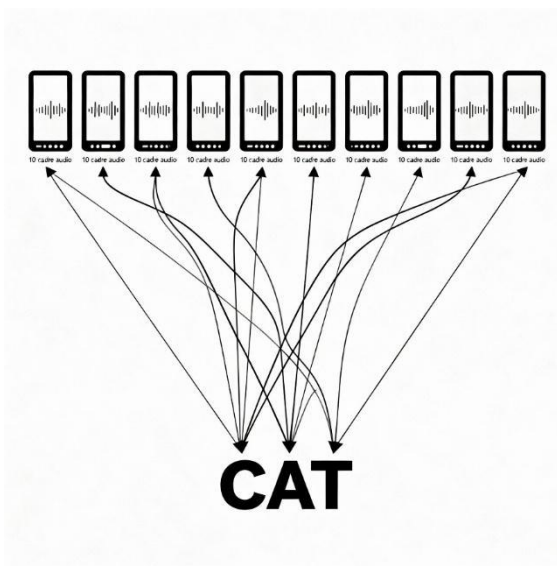


Fig.2 Connectionist Temporal Classification

Mecanismul CTC (fig. 2) se bazează pe câteva concepte cheie:

1. *Simbolul "Blank" (ϵ):* CTC extinde setul de etichete de ieșire (ex: alfabetul) cu un simbol

special, numit "blank". Acest simbol permite modelului să nu emită nicio etichetă la un anumit pas de timp, acționând ca un separator sau un "spațiu de așteptare".

2. *Căi de Aliniere și Colapsare*: Rețeaua neuronală produce o distribuție de probabilitate peste setul extins de etichete (alfabet + blank) pentru fiecare cadru audio de intrare. O "cale" (path) este o secvență de etichete de lungime T. CTC definește o regulă de colapsare pentru a mapa aceste căi la transcrierea finală:

- Mai întâi, se elimină toate simbolurile "blank".
- Apoi, se elimină caracterele repetate consecutive.

De exemplu, căile c-aa-t, cc-a-t- și -ca-tt vor fi toate colapsate la transcrierea finală "cat".

3. *Calculul Funcției de Pierdere*: Probabilitatea unei transcrieri țintă este calculată ca suma probabilităților tuturor căilor valide care se colapsează la acea transcriere. O abordare naivă ar fi computațional explozivă. CTC utilizează un algoritm eficient de programare dinamică, similar cu algoritmul Forward-Backward din HMM-uri, pentru a calcula această sumă. Funcția de pierdere CTC este apoi definită ca log-verosimilitatea negativă a transcrierii corecte.

Deși este un mecanism puternic, CTC are o limitare importantă: face o ipoteză de independență condițională, presupunând că predicțiile la fiecare pas de timp sunt independente, dat fiind întregul semnal de intrare. Acest lucru face dificilă modelarea dependențelor lingvistice pe termen lung direct la nivelul funcției de cost.

5.4. Arhitecturi Encoder-Decoder cu Atenție

Pentru a depăși limitarea de independență a CTC, o a doua familie de modele end-to-end s-a inspirat din arhitecturile de tip encoder-decoder utilizate cu succes în traducerea automată neuronală. Modelul emblematic pentru ASR din această categorie este Listen, Attend and Spell (LAS), propus de Chan et al. (2016).

O arhitectură LAS este compusă din trei module principale:

1. Encoder ("Listen"): Un encoder, de obicei un RNN bidirecțional (Bi-LSTM) sau, mai recent, un Transformer, procesează întreaga secvență de trăsături acustice de intrare și o transformă într-o secvență de reprezentări de nivel înalt, bogate în context.

2. Mecanism de Atenție ("Attend"): Acesta este elementul cheie. La fiecare pas de generare a ieșirii, mecanismul de atenție calculează o distribuție de ponderi peste reprezentările de ieșire ale encoderului. Aceste ponderi indică pe ce parte a semnalului audio ar trebui să se "concentreze" modelul pentru a genera următorul caracter.

3. Decoder ("Spell"): Un decoder, de obicei un RNN unidirecțional, generează secvența de caractere de ieșire, unul câte unul. La fiecare pas, decoderul primește caracterul generat anterior și un vector de context, care este o sumă ponderată a ieșirilor encoderului, calculată pe baza ponderilor de atenție.

Mecanismul de atenție nu doar că îmbunătățește performanța, permițând modelului să gestioneze dependențe pe termen lung, dar oferă și un grad sporit de interpretabilitate: prin vizualizarea matricelor de atenție, putem observa cum se aliniază generarea textului cu semnalul

audio.

5.5. Modele Hibride: CTC + Atenție

Într-un efort de a combina avantajele ambelor paradigme, cercetătorii au propus modele hibride care utilizează simultan atât o funcție de pierdere CTC, cât și un mecanism de atenție în timpul antrenării. În această configurație, funcția de pierdere totală este o combinație liniară a celor două:

$$L_{total} = \lambda \cdot L_{CTC} + (1 - \lambda) \cdot L_{Attention}$$

Această abordare hibridă s-a dovedit extrem de eficientă. Funcția CTC, care forțează o aliniere monotonică, acționează ca un regularizator și un ghid pentru mecanismul de atenție, accelerând convergența și crescând stabilitatea, în special în etapele inițiale ale antrenării. Pe de altă parte, componenta de atenție permite modelului să învețe dependențe lingvistice mai complexe. Miao et al. (2020) au demonstrat că această arhitectură este deosebit de eficientă pentru sistemele online, unde latența este un factor critic.

5.6. Învățarea Auto-Supervizată: Wav2Vec 2.0

O revoluție recentă în ASR este reprezentată de introducerea învățării auto-supervizate (self-supervised learning - SSL), în special prin modelele Wav2Vec și Wav2Vec 2.0, dezvoltate de Facebook AI (acum Meta AI). Aceste modele abordează una dintre cele mai mari provocări din domeniu: dependența de cantități masive de date audio *etichetate*.

SSL permite pre-antrenarea unui model puternic pe sute sau mii de ore de audio neetichetat, învățând reprezentări acustice generale. Modelul Wav2Vec 2.0, propus de Baevski et al. (2020), are o arhitectură formată dintr-un encoder convoluțional și un encoder Transformer.

Procesul de antrenare se desfășoară în două etape:

1. Pre-antrenare (Auto-supervizată): Modelul este antrenat pe audio neetichetat. O parte din secvența de intrare este mascată aleatoriu, iar sarcina modelului este de a prezice o reprezentare cuantizată a segmentelor mascate, pe baza contextului furnizat de segmentele nemascate. Această sarcină forțează modelul să învețe reprezentări bogate și contextualizate ale semnalului vocal.
2. Fine-tuning (Supervizat): După pre-antrenare, un strat de clasificare liniar este adăugat deasupra encoderului Transformer, iar întregul model este ajustat (fine-tuned) pe o cantitate mult mai mică de date etichetate, folosind de obicei funcția de pierdere CTC.

Impactul acestei abordări este major: a demonstrat că este posibil să se obțină performanțe de vârf cu de până la 100 de ori mai puține date etichetate decât abordările complet supervizate, democratizând astfel dezvoltarea de sisteme ASR pentru limbile cu resurse limitate.

5.7. Recunoașterea Vocii în Context Multilingv și Low-Resource

Recunoașterea automată a vorbirii a atins performanțe remarcabile în limbile cu resurse bogate, dar rămâne adesea ineficientă în contexte cu limbi minoritare, dialecte, sau în scenarii multilingve reale. Deși limbi precum engleza, chineza mandarină sau germana beneficiază de corpuri vaste de date audio și resurse lingvistice (dicționare fonetice, transcrieri), marea majoritate a limbilor

lumii sunt considerate „low-resource”, adesea din lipsă de infrastructură digitală sau din cauza neglijării acestora în cercetare.

În această etapă a cursului, vom analiza provocările conceptuale și tehnice asociate contextului multilingv, prezentând strategii moderne precum modele multilingve, învățare auto-supervizată și antrenare transferabilă. Vom discuta, de asemenea, aspectele culturale și etice implicate în proiectarea ASR pentru limbi minoritare.

5.7.1. Diversitate lingvistică și provocări fonologice în ASR multilingv

Recunoașterea automată a vorbirii (ASR) în contexte multilingve presupune adaptarea la o diversitate fonetică și fonologică ridicată. Fiecare limbă prezintă particularități acustice care pot compromite performanța modelelor generaliste. De exemplu, sistemele tonale, precum vietnameza sau chineza mandarină, utilizează variații de înălțime fonemică pentru a distinge sensurile cuvintelor. Alte limbi, precum georgiana, conțin consoane ejective, în timp ce limbi bantu precum xhosa sau zulu prezintă click-uri fonemice — sunete articulatorii absente în limbile indoeuropene. Modelele ASR antrenate predominant pe limbi occidentale, cum sunt engleza sau germana, tind să nu captureze aceste fenomene acustice specifice, ceea ce duce la creșteri semnificative ale ratei de eroare în recunoaștere.

De asemenea, morfologia unei limbi influențează direct granularitatea unităților recunoscute. Limbile aglutinative, cum ar fi finlandeza sau turca, generează un număr extrem de mare de forme cuvânt prin concatenarea morfemelor, fapt ce afectează negativ eficiența modelelor bazate pe recunoașterea cuvintelor întregi (word-based ASR). Modele ASR dezvoltate pentru limbi izolante, precum engleza, eșuează sistematic atunci când sunt aplicate pe limbi polisintetice, unde informația sintactică și semantică este comprimată în unități morfologice complexe. Astfel, devine stringentă necesitatea dezvoltării unor reprezentări acustice mai generale, independente de particularitățile limbii, care să capteze trăsături articulatorii universale.

O provocare deosebită, adesea neglijată în literatura de specialitate, este *code-switching*, adică alternarea între două sau mai multe limbi în cadrul aceleiași propoziții sau interacțiuni discursive. Acest fenomen este omniprezent în societăți plurilingve, fiind documentat frecvent în India (hindi–engleză), Mexic (spaniolă–nahuatl), Singapore (mandarină–engleză), dar și în România (română–maghiară sau romani). Modelele ASR clasice, antrenate sub ipoteza unei limbi unice la intrare, se confruntă cu dificultăți semnificative atunci când trebuie să recunoască astfel de secvențe amestecate. Schimbările bruște de fonologie și sintaxă, precum și suprapunerea lexicului între limbi, complică segmentarea și interpretarea corectă a vorbirii. Conform lucrării lui Khassanov et al. (2022), abordarea eficientă a code-switching-ului presupune construirea de vocabular multilingv, partajarea parametrilor acustici între limbi și utilizarea unor mecanisme dinamice de detectare a limbii active. Modelul propus de aceștia, bazat pe rețele LSTM cu atenție (attention-based LSTM), integrează identificarea limbii și transcrierea într-un flux unificat de inferență, demonstrând o îmbunătățire semnificativă în scenarii code-switched reale.

Antrenarea simultană a unui model pe mai multe limbi permite învățarea unor reprezentări acustice partajate, care pot fi ulterior adaptate la limbi necunoscute prin fine-tuning. Un exemplu relevant este modelul Wav2Vec 2.0 XLSR, propus de Baevski et al. (2021), care utilizează un encoder

comun pentru reprezentarea fonetică a peste 50 de limbi, urmat de straturi de ieșire specializate pentru fiecare limbă în parte. Performanțele acestui model sunt impresionante: pe limbile africane amharică și igbo, rata de eroare la recunoaștere (WER) a fost redusă cu până la 54% comparativ cu modelele monolingve. De asemenea, modelul permite transferul zero-shot — adică recunoașterea în limbi neîntâlnite în timpul antrenării — cu rezultate semnificative în swahili și vietnameză. Este remarcabil faptul că, după doar o oră de fine-tuning cu date transcrise, modelul depășește performanța modelelor antrenate from-scratch pe peste 100 de ore de date (Baevski et al., 2021).

Un alt exemplu emblematic este modelul Whisper, dezvoltat de OpenAI și prezentat de Radford et al. (2022), antrenat pe un corpus masiv de peste 680.000 de ore de vorbire în 96 de limbi. Whisper este capabil să recunoască vorbirea în condiții acustice dificile, să detecteze automat limba activă și chiar să realizeze traducere automată direct din semnalul vocal. Acest model deschide perspective noi pentru ASR în regiuni cu infrastructură limitată și în contexte medicale sau educaționale.

Mai mult, recentul model IPA-Speech propus de Tong et al. (2023) introduce o paradigmă alternativă în codificarea fonetică. În loc de utilizarea fonemelor definite la nivel de limbă, IPA-Speech propune reprezentarea fonetică prin trăsături articulatorii universale, codificate conform alfabetului fonetic internațional (IPA). Această strategie permite modelelor să generalizeze mai bine pe limbi necunoscute sau slab reprezentate. De exemplu, trăsătura [voiced bilabial plosive], care corespunde fonemului /b/, este prezentă în peste 100 de limbi, inclusiv în română, swahili sau tagalog. Prin învățarea unor astfel de trăsături independente de limbă, modelele pot obține o formă de abstracție fonologică supralingvistică, extrem de utilă în scenarii multilingve.

5.7.2. ASR pentru limbi low-resource

În paradigma clasică ASR, dezvoltarea unui sistem robust presupune existența a zeci sau chiar sute de ore de înregistrări vocale transcrise, dicționare fonetice detaliate și modele lingvistice elaborate. Din păcate, pentru majoritatea limbilor minoritare sau indigene, aceste resurse lipsesc sau sunt extrem de limitate. Spre exemplu, limba Mapudungun, vorbită în Chile, dispune de mai puțin de 2000 de ore de înregistrări disponibile public, dintre care foarte puține sunt însoțite de transcrieri.

În acest context, modelele bazate pe învățare auto-supervizată (SSL) oferă o alternativă viabilă. Arhitecturi precum HuBERT, Wav2Vec 2.0 și UniSpeech sunt capabile să extragă reprezentări fonetice din date audio brute, fără a necesita etichetare prealabilă. Aceste modele identifică regularități acustice prin metode de clustering latent și pot fi ulterior adaptate la limbile țintă prin fine-tuning cu doar câteva minute de date etichetate. Conform benchmark-ului SUPERB realizat de Yang et al. (2021), un model Wav2Vec 2.0, fine-tuned pe mai puțin de 10 minute de date, atinge performanțe comparabile cu modele tradiționale antrenate pe peste 100 de ore de date complet transcrise. Această capacitate de adaptare rapidă transformă modelele SSL într-o soluție eficientă pentru dezvoltarea ASR în limbi rare sau slab documentate.

5.7.3. Studii de caz și inițiative comunitare

Mozilla Common Voice: democratizarea colectării de date vocale

Mozilla Common Voice (Ardila et al. 2019) este una dintre cele mai importante inițiative open-

source dedicate dezvoltării corpusurilor vocale multilingve, în special pentru limbile slab reprezentate în spațiul digital. Proiectul își propune să sprijine construirea de modele de recunoaștere automată a vorbirii (ASR) prin colectarea participativă de date vocale etichetate. Până în 2024, platforma susținea peste 120 de limbi, inclusiv limbi rare sau în pericol, cum ar fi galeza, zhuang sau bashkir .

Aceste date sunt colectate de către voluntari din comunități locale, printr-un proces deschis și transparent, în care participanții își exprimă explicit consimțământul, în conformitate cu licențierea Creative Commons (CC-BY 4.0). Conform unui studiu realizat de Oluremi et al. (2025), Common Voice a atins un volum de peste 23 de milioane de înregistrări vocale, creând premisele pentru antrenarea modelelor de ASR în limbi cu resurse reduse sau inexistente anterior .

Common Voice nu este doar o platformă tehnologică, ci și un exemplu de infrastructură etică și participativă în AI lingvistic, permițând comunităților să controleze modul în care datele lor vocale sunt colectate și folosite. În colaborare cu organizații precum Te Hiku Media (Noua Zeelandă), Mozilla sprijină și inițiative indigene pentru păstrarea limbilor minoritare prin tehnologie vocală .

Masakhane ASR: construire colaborativă pentru limbile africane

Masakhane ASR este o inițiativă continentală lansată în Africa pentru a aborda excluderea limbilor africane din tehnologiile moderne de vorbire. Masakhane promovează o abordare colaborativă, în care cercetători, lingviști, dezvoltatori și membri ai comunităților locale contribuie împreună la construirea sistemelor ASR.

Metodologia implică selecția locală a datelor — inclusiv interviuri, conversații din viața cotidiană și corpusuri jurnalistice — și fine-tuning-ul modelelor existente (precum Wav2Vec 2.0 sau XLS-R) pe aceste date. Accentul se pune pe învățarea echitabilă și controlul comunitar asupra resurselor lingvistice.

Rezultatele sunt deja tangibile: conform Blocker et al. (2025), modelele Masakhane antrenate pentru limbile isiXhosa, Wolof și Luganda obțin performanțe funcționale (WER acceptabil) în aplicații reale precum mesageria vocală, interviurile asistate sau sistemele de răspuns vocal automatizat . Inițiativa este recunoscută nu doar pentru inovația tehnologică, ci și pentru valorile de incluziune și justiție lingvistică pe care le promovează.

Whisper în aplicații medicale: un exemplu de utilitate socială

Modelul Whisper, dezvoltat de OpenAI, s-a dovedit deosebit de robust în contextul limbilor low-resource, datorită capacității sale de a generaliza pe date diverse și zgomotoase. Fine-tuning-ul lui Whisper pe limbi indigene și context medical a deschis noi frontiere în domeniul sănătății publice, în special în regiunile subfinanțate.

Într-un exemplu recent documentat de Blocker et al. (2025), Whisper a fost utilizat în spitale rurale din India pentru transcrierea automată a consultațiilor în limbi precum bhojpuri și maithili — limbi cu puține resurse digitale și lexic medical slab standardizat. Sistemul a fost folosit nu doar pentru transcriere, ci și pentru extragerea automată de entități medicale (nume de boli, simptome, tratamente), facilitând astfel documentarea și accesul la informație în medii unde personalul medical este suprasolicitat sau lipsit de infrastructură digitală .

De asemenea, în combinație cu sisteme TTS (Text-to-Speech), Whisper a fost integrat în aplicații de sprijin pentru pacienți cu dizabilități de exprimare — precum cei afectați de dizartrie severă — permițând conversii bidirecționale între voce sintetizată și recunoaștere vocală în limba maternă.

Aceste exemple nu doar că demonstrează aplicabilitatea modelelor ASR în contexte critice, dar evidențiază și importanța unui design etic, localizat și centrat pe utilizator.

Mozilla Common Voice, Masakhane ASR și utilizarea Whisper în domeniul medical nu reprezintă doar inovații tehnologice, ci inițiative care întruchipează valori de participare comunitară, respect cultural și justiție digitală. Într-o epocă în care AI-ul lingvistic riscă să amplifice inegalitățile, aceste proiecte oferă contraexemple constructive, demonstrând că infrastructura tehnologică poate fi proiectată cu și pentru comunități.

Tehnologia vocală nu este un demers neutru din punct de vedere cultural sau politic. Selecția limbilor suportate, preferința pentru anumite variante dialectale și standardizarea pronunției reflectă alegeri implicite care, deși motivate tehnic, pot avea consecințe semnificative asupra diversității lingvistice. De exemplu, preferința pentru forme standardizate, precum euskara batua în locul dialectelor regionale basce, poate duce la marginalizarea variantelor locale și la erodarea patrimoniului oral autentic.

Un alt aspect problematic îl reprezintă practicile de colectare a datelor vocale fără consimțământ informat. Există cazuri documentate în care asistenți vocali precum Amazon Alexa au stocat și utilizat mostre vocale fără o notificare clară către utilizatori, ridicând probleme de transparență și etică a datelor biometrice.

Mai mult, în unele inițiative comerciale, vocea colectată din comunități este reutilizată în produse generative sau sisteme TTS fără ca beneficiile economice sau accesul tehnologic să fie redistribuite către acele comunități. Acest dezechilibru între sursă și beneficiu accentuează inegalitățile deja existente.

Ca răspuns la aceste riscuri, comunități și organizații internaționale au început să dezvolte mecanisme de protecție etică. Platforme precum Mozilla Common Voice aplică licențiere deschisă (CC-BY) și necesită consimțământ explicit în procesul de donare a vocii. Inițiativa africană Masakhane integrează audit comunitar activ, asigurându-se că datele sunt colectate, procesate și folosite în mod transparent și în beneficiul direct al vorbitorilor. De asemenea, organizații precum Linguistic Data Consortium (LDC) și European Language Resources Association (ELRA) propun ghiduri etice pentru documentarea, partajarea și utilizarea responsabilă a resurselor lingvistice.

În această lumină, dezvoltarea de sisteme ASR multilingve pentru limbi cu resurse limitate trebuie privită nu doar ca o provocare tehnologică, ci și ca un act de justiție lingvistică. Modele de scară largă precum Whisper sau Wav2Vec XLSR oferă un cadru promițător pentru incluziune, dar succesul real al acestor tehnologii depinde de integrarea unor principii suplimentare. Printre acestea se numără participarea activă a comunităților locale în procesul de colectare și validare a datelor, stabilirea unor standarde clare de consimțământ și protecție a datelor vocale și, nu în ultimul rând, recunoașterea explicită a valorii culturale și lingvistice pe care o reprezintă fiecare limbă sau dialect

în parte.