



Marta Andrić

Filozofski fakultet Sveučilišta u Zagrebu

mandric2@m.ffzg.unizg.hr

Tatjana Podvezanec

Čakovec, Hrvatska

podvezanec@gmail.com

Glagolski oblici u turskim televizijskim serijama: korpusna analiza kao resurs za istraživanje nenamjernoga usvajanja jezika

Cilj ovoga rada jest dokumentirati i barem djelomice opisati trend usvajanja turskoga jezika putem televizijskih serija u Hrvatskoj – fenomen prisutan od 2010. godine, kad je na jednoj domaćoj televiziji prvi put emitirana turska serija. Od tad je praćenje turskih serija, dodatno potaknuto internetskom dostupnošću sadržaja, postalo stabilan i izrazito živ oblik kontakta hrvatskih gledatelja s turskim jezikom.

U radu se analizira zastupljenost glagolskih oblika u prvih pet epizoda turske televizijske serije *Yabani* 'Divlji' kako bi se procijenilo do koje razine gledatelji serija mogu spontano usvajati turski jezik. Korpus od 32 253 riječi obrađen je u alatu *Sketch Engine*, uz nužna ručna ispravljanja zbog nedostatne morfološke anotacije za turski jezik. Usporedba s rječnikom *A Frequency Dictionary of Turkish* pokazala je da se oko 50 % najfrekventnijih lema podudara s jezikom serije, dok udio podudaranja najčešćih glagola iznosi 74 %.

Analizom korpusa izdvojeni su svi glagolski oblici, izračunate njihove frekvencije te utvrđeni međusobni omjeri i distribucije. Najbrojniji oblici u korpusu su prezent glagola *imek* 'biti', imperativ i perfekt na *-DI*, dok su infinitivne forme i složena vremena znatno rjeđa. Analiza glagolskih kategorija pokazuje da su u korpusu najbrojniji oblici u jednini: 1. i 2. lice zbog dijaloške strukture teksta, a 3. lice zbog narativnih dijelova.

Istraživanje provedeno s gledateljima serija (N = 15) pokazuje da ispitanici najuspješnije prepoznaju finitne oblike poput prezenta i imperativa, dok je uspješnost kod infinitivnih i složenijih oblika osjetno niža. Time su potvrđene početne pretpostavke da se gledanjem serija najlakše usvajaju osnovne finitne glagolske strukture.

Ključne riječi: turski jezik; televizijske serije; korpusna analiza; glagolski oblici; nenamjerno usvajanje jezika

1. Uvod

Već je petnaest godina kako se na televizijskim programima u Hrvatskoj prikazuju turske dramske serije, tzv. *sapunice*. S velikim interesom gledatelja za te serije porastao je i interes za turski jezik i kulturu. U skladu s time, među vrlo širokom gledateljskom populacijom već je godinama zamjetno poznavanje turskoga jezika na raznim razinama: od minimalnoga znanja koje se svodi na nekoliko riječi i fraza, do znatno naprednijih razina.

Cilj ovoga rada bio je korpusnom analizom jezika turskih televizijskih serija istražiti koje se razine jezičnoga znanja mogu usvojiti njihovim praćenjem. U tu je svrhu izrađen korpus koji se sastoji od prvih pet epizoda serije *Yabani* ('Divlji', Turska, 2023).¹ Iz toga korpusa izdvojeni su najčešći glagolski oblici koji su potom razvrstani prema referentnim razinama poznavanja jezika opisanim u Zajedničkome europskom referentnom okviru za jezike (ZEROJ). Kao fokus istraživanja izabrani su upravo glagolski oblici jer oni predstavljaju temeljni strukturni element rečenice, ali i zato što omogućuju relativno jednostavno i pouzdano praćenje razina jezičnoga znanja koje gledatelji mogu usvojiti. Rezultati dobiveni korpusnom analizom poslužili su kao osnova za izradu upitnika za provjeru poznavanja tih oblika kod ispitanika izloženih turskom jeziku putem serija. Na temelju toga upitnika provedeno je istraživanje među studentima i polaznicima tečaja turskoga jezika koji su se nalazili na samom početku formalnoga obrazovanja, ali su imali predznanje turskoga jezika koje su stekli upravo praćenjem serija. Tako je omogućena usporedba rezultata korpusne analize, frekvencije i distribucije pojedinih glagolskih oblika, te njihove prepoznatljivosti i usvojenosti kod gledatelja izloženih turskomu jeziku putem audiovizualnih sadržaja.

2. Teorijski okvir

Jezik televizijskih serija, kao dio širega područja jezika audiovizualnih sadržaja, obrađivan je u brojnim radovima, no istraživanja na građi turskoga jezika razmjerno su malobrojna. Fenomenu turskih televizijskih serija posvećen je određen broj radova i doktorskih disertacija koji pripadaju drugim istraživačkim područjima. Primjerice, doktorska disertacija Şerifa Arslana *Yumuşak güç algısında Türk televizyon dizileri* („Turske televizijske serije u kontekstu percepcije meke moći”) (2022) nastala je u okviru sociologije. Faruk Sadıç (2024) u članku *Azerbaijani TV Viewers' Practices of Watching Turkish Television Programs in the Context of Cultural*

1 Ovaj rad polazi od diplomskoga rada Tatjane Podvezanec obranjenog 2025. godine na Filozofskom fakultetu Sveučilišta u Zagrebu. Naslov diplomskoga rada je *Računalna analiza i minimalni vokabular turskih dramskih serija na primjeru serije Yabani*. Oba rada imaju isti korpus, ali je njihova sadržajna podudarnost minimalna. Diplomski je rad usmjeren na leksik i minimalni vokabular navedene serije, a ovaj se rad bavi analizom glagolskih oblika, njihovom distribucijom u korpusu te njihovom ulogom u nenamjernom usvajanju turskoga jezika. Autorice zahvaljuju prof. dr. sc. Ivani Franić na vrijednim uredničkim sugestijama tijekom pripreme rada za objavu.

Studies: Turkish Culture and Tourism fenomen serija sagledava unutar kulturalnih studija, dok istraživanje Elif Nuroğlu (2013) o utjecaju turskih serija na „serijski turizam“ i turističke tokove na Bliskom istoku i Balkanu pripada području ekonomije.

Dosadašnja istraživanja jezika turskih audiovizualnih medija uglavnom su usmjerena na filmove, dok su televizijske serije rjeđe predmet analize (Eryılmaz, 2024, str. 889). Pritom se audiovizualni sadržaji pretežno promatraju iz perspektive poučavanja turskoga kao stranoga jezika i njihove primjene u nastavnome procesu, a ne iz aspekta nenamjernoga usvajanja jezika. Eryılmaz (2024) također upozorava na metodološke nedostatke i manjak preciznosti u oblikovanju dosadašnjih istraživanja. Metodološke manjkavosti dosadašnjih istraživanja povezane su i s činjenicom da korpusnolingvistički pristup jeziku televizijskih serija u turskome kontekstu još uvijek nije dovoljno razvijen. Pojedini radovi primjenjuju korpusnolingvistički pristup u analizi jezika televizijskih serija, ali ne na građi turskoga jezika. Primjerice, rad Hatice Sezgin i Mustafe Serkana Öztürka (2020) *A corpus analysis on the language on TV series* bavi se sličnim pitanjem kao naše istraživanje, premda je usmjeren isključivo na vokabular u serijama na engleskom jeziku, pri čemu autori kao izvor koriste *British TV Series Corpus* (BTSC) (Sezgin i Öztürk, 2020, str. 238). Koliko je poznato, ne postoje radovi koji korpusnolingvističku analizu jezika turskih televizijskih serija povezuju s istraživanjem nenamjernoga usvajanja jezika, kako je to učinjeno u našem radu.

Izvan granica Turske i turskoga jezika istraživanja posvećena jeziku televizijskih serija obuhvaćaju različite teorijske i metodološke pristupe i velikim su dijelom usmjerena na serije na engleskome i španjolskome jeziku. Ta se istraživanja najčešće smještaju u okvire primijenjene lingvistike, korpusne lingvistike, lingvistike medija i sociolingvistike, pri čemu se navedeni pristupi nerijetko i međusobno kombiniraju (Bednarek, 2018, str. 17). U ovome je radu naglasak stavljen na korpusnu analizu, a dobiveni se rezultati dodatno interpretiraju i iz aspekta nenamjernoga usvajanja jezika. Na taj se način jezične distribucije i frekvencijski obrasci analizirani u korpusu dovode u vezu s mogućnostima njihova spontanog usvajanja tijekom izlaganja jeziku medija.

Korpusna lingvistika predstavlja važan metodološki pristup u istraživanju turskoga jezika. Danas su dostupni nacionalno konstruirani, uravnoteženi korpusi, poput *Turkish National Corpus* koji obuhvaća približno 50 milijuna riječi suvremenoga turskog jezika i temelji se na kontroliranom uzorku različitih tekstnih vrsta (Aksan i sur., 2012), kao i raniji METU korpus (*Middle East Technical University Corpus*), prvi uravnoteženi korpus turskoga jezika koji sadržava približno 1,7 milijuna pojava iz oko tisuću pisanih tekstova različitih žanrova, poput romana, eseja, znanstvenih članaka, intervjua i novinskih tekstova (Çöltekin i sur., 2023, str. 451). Uz njih postoje i veliki web-korpusi kao što je, primjerice, *trTenTen*, dostupan na platformi *Sketch Engine*, koji sadrži oko 4,9 milijardi riječi, ali je automatski prikupljen s interneta te stoga žanrovski i stilski manje kontroliran.² Pregled postoje-

2 trTenTen – Turkish corpus from the web | Sketch Engine, 5.5.2026.

ćih korpusnih resursa pokazuje da su za turski jezik razvijeni brojni opći, anotirani i specijalizirani korpusi (Çöltekin i sur., 2023, str. 451–465). No, ne postoji korpus jezika turskih televizijskih serija, pa u tom smislu ovo istraživanje, makar i malenoga opsega, predstavlja novost.

Istraživanja nenamjernoga usvajanja jezika putem audiovizualnih sadržaja polaze od pretpostavke da se jezični elementi mogu usvajati tijekom aktivnosti usmjerenih na razumijevanje sadržaja, a ne na svjesno učenje jezika (d'Ydewalle i Pavakanun, 1995; Teng, 2024). U literaturi se pritom ističe da gledanje televizijskih programa i serija može pridonijeti usvajanju vokabulara, ali i drugih jezičnih elemenata, uključujući morfološke obrasce, kolokacije i obrasce uporabe jezika (Teng, 2024, str. 2–6). Posebno se naglašava uloga vizualnoga konteksta i titlova koji mogu olakšati razumijevanje i uočavanje jezičnih oblika u govornome *inputu* (Montero Perez, Sydorenko i Huntley, 2024, str. 40–41). Radovi posvećeni nenamjernomu usvajanju turskoga jezika putem televizijskih serija ili drugih audiovizualnih sadržaja, koliko je poznato, zasad nisu zabilježeni.

3. Metodološki okvir

Kao što je prethodno navedeno, korpus istraživanja čini prvih pet epizoda turske televizijske serije *Yabani*.³ Ta je serija odabrana zato što je u vrijeme planiranja istraživanja bila aktualna i dostupna hrvatskoj publici. I druge bi serije mogle biti prikladan korpus za istraživanja ove vrste. Jednako tako, i uključivanje epizoda iz više serija te proširenje korpusa vjerojatno bi omogućilo dodatne uvide ili potvrdilo dobivene rezultate. No, s obzirom na to da je riječ o području koje još nije opsežnije istraženo, odabrani korpus smatran je prikladnim i dostatnim za potrebe analize.

U radu su postavljena dva temeljna istraživačka pitanja. Prvo se odnosi na korpusnu analizu jezika televizijske serije *Yabani* te je usmjereno na pitanje kakvi su distribucija i čestotnost glagolskih oblika u analiziranom korpusu. U okviru toga pitanja analizirani su svi glagolski oblici izdvojeni iz korpusa, njihova frekvencija, međusobni omjeri te raspodjela prema pojedinim glagolskim kategorijama, uključujući finitne i infinitne oblike te distribuciju prema licu i broju. Cilj je bio utvrditi u kojim su omjerima glagolske strukture zastupljene u jeziku analizirane serije te koje se gramatičke kategorije zbog svoje frekventnosti mogu smatrati posebno relevantnima u kontekstu usvajanja jezika putem televizijskih sadržaja.

Drugo istraživačko pitanje usmjereno je na prepoznavanje i usvajanje glagolskih oblika među gledateljima izloženima turskom jeziku putem televizijskih serija. U tome se dijelu rada nastojalo utvrditi koji se glagolski oblici najuspješnije

3 Serija *Yabani* počela se u Turskoj prikazivati 12. rujna 2023. na kanalu NOW. Redatelj je Çağatay Tosun, a scenarij su napisale Hilal Yıldız i Yekta Torun. Radnja prati Yamana, mladića koji nakon odrastanja na ulici otkriva da zapravo potječe iz jedne imućne obitelji. Takvim se sadržajem u seriji uspostavlja kontrast dvaju društvenih slojeva i njima pripadajućih jezičnih registara. U Hrvatskoj se serija počela prikazivati u listopadu 2024. na Novoj TV pod naslovom *Divlje srce*.

prepoznaju i usvajaju te postoji li povezanost između njihove frekvencije u korpusu i njihove prepoznatljivosti među ispitanicima.

Uz navedena istraživačka pitanja tijekom rada otvorio se i niz metodoloških pitanja povezanih s obradom korpusa i mogućnostima njegove računalne analize. Poseban je problem predstavljala činjenica da korpusni alat *Sketch Engine*⁴ (Kilgarrieff i sur., 2014) koji je korišten za analizu ne omogućuje potpunu i pouzdanu automatsku morfološku obradu turskoga jezika, zbog čega je dio analize bilo potrebno provesti ručno. To je ujedno utjecalo i na opseg korpusa koji je morao ostati dovoljno velik da omogući reprezentativne rezultate, ali i dovoljno ograničen da analiza bude metodološki provediva u okviru ovoga istraživanja.

U korpusnoj lingvistici vrijedi pravilo da veći korpusi daju stabilnije rezultate, no u području specijaliziranih korpusa izvori navode da se analitički relevantni podaci mogu dobiti već iz skupa od oko 30 000 riječi (Bowker i Pearson, 2002, str. 144). Korpus analiziran u ovome radu sastavljen je od transkripata pet epizoda serije koji su objedinjeni u jedinstven tekstualni skup. Obuhvaća 32 253 riječi (odnosno pojava). Računalnom obradom korpusa pojava u korpusnom alatu *Sketch Engine* dobiven je frekvencijski popis od 7007 različenica.

Za obradu turskoga jezika, osobito u području morfološke analize, razvijen je niz specijaliziranih alata kao što su sustav *Zemberek*⁵ te noviji alati poput *ITUNLP pipeline*⁶ i dr. (Çoltekin i sur., 2023). Ti su alati učinkoviti u analizi pojedinačnih oblika i njihovih morfoloških struktura. Međutim, njihova je primarna namjena analiza morfološke strukture i gramatičkih obilježja pojedinačnih riječi, dok su mogućnosti rada s velikim korpusima i fleksibilnog pretraživanja ograničene. U ovom je radu stoga korišten alat *Sketch Engine* koji omogućuje izgradnju i pretraživanje velikih korpusa te formuliranje upita u CQL jeziku. Unatoč ograničenjima u morfolškoj anotaciji za turski jezik, njegova funkcionalnost za korpusnu analizu pokazala se prikladnijom za ciljeve ovoga istraživanja.

Turski jezik pripada skupini aglutinativnih jezika u kojima se gramatičke kategorije izražavaju nizovima sufikasa koji se linearno nadovezuju na osnovu riječi. Jedan površinski oblik može sadržavati veći broj morfema i odgovarati složenijim sintaktičkim strukturama u flektivnim jezicima. Takva morfolška složenost dovodi do visoke razine morfološke dvoznačnosti jer isti niz fonema može imati više mogućih morfoloških analiza, ovisno o kontekstu (Ofłazer i Saraçlar, 2018, str. 1–10). Iako su za turski jezik razvijeni sustavi za automatsku morfolšku anotaciju, njihova primjena u korpusnim istraživanjima često je ograničena upravo zbog te morfološke dvoznačnosti.

Takva se ograničenja odražavaju i na mogućnosti preciznoga korpusnog pretraživanja u alatu *Sketch Engine*. Lematizacija u tome sustavu nije moguća, a morfolška anotacija nije dovoljno diferencirana da bi omogućila precizno razlučiva-

4 Dostupno na: <https://www.sketchengine.eu/>

5 <https://github.com/ahmetaa/zemberek-nlp>, 10.4.2026.

6 ITU Turkish Natural Language Processing Web Interface, 10.4.2026.

nje gramatičkih funkcija. Primjerice, ako se želi istražiti sve slučajeve korištenja sufiksa *-(y)Im*, 1. lica prezenta glagola *imek* 'biti', zbog srodnosti i homografije toga sufiksa sa sufiksima za tvorbu raznih drugih gramatičkih oblika, CQL upitom *[tag="*.cpl:pres.*<1s>"]* dobivaju se rezultati u kojima se pojavljuju ne samo 1. lice jednine prezenta glagola *imek*, nego i oblici 1. lica jednine svih jednostavnih i složenih glagolskih vremena, potom optativa, necesitativa, posibilitativa i imposibilitativa, zatim proparticipa na *-DIk* i *-(y)AcAk*, sve imenice na kojima je izražena posvojnost 1. lica jednine, kao i sve pojavnice čija linearna fonološka struktura završava jednakom sekvencom (npr. zamjenica *kim*, imenica *yadım* i sl.), i svi su ti oblici u sustavu označeni kao da su glagol (*<V>cpl:pres<1s>*). Zbog ovakvih ograničenja veliki dio obrade korpusa bilo je potrebno ručno provesti ili barem provjeriti.

Drugi dio istraživanja odnosio se na provjeru prepoznatljivosti i usvojenosti glagolskih oblika među gledateljima turskih serija. U tu je svrhu sastavljen upitnik temeljen na glagolskim oblicima izdvojenima korpusnom analizom te raspoređeni prema referentnim razinama poznavanja jezika opisanima u Zajedničkome europskom referentnom okviru za jezike. Istraživanje je provedeno među studentima i polaznicima tečaja turskoga jezika koji su bili izloženi turskom jeziku putem televizijskih serija, a cilj mu je bio usporediti rezultate korpusne analize s prepoznatljivošću pojedinih glagolskih oblika među ispitanicima. Budući da je riječ o metodološki zasebnom segmentu rada, detaljan opis istraživanja, uzorka i strukture upitnika donosi se u poglavlju 6. *Gledanje serija kao izvor inputa: istraživanje i analiza.*

4. Čestotnost leksika u jeziku serije i usporedba s referentnim korpusom

Jedno od pitanja koje se postavilo na početku ovoga istraživanja bilo je pitanje kakav je odnos između jezika serije i standardnog turskog i koliko je, s obzirom na to, jezik televizijskih serija prikladan kao medij za usvajanje i učenje turskoga jezika. Radi dobivanja okvirnoga uvida u to pitanje uspoređena je učestalost leksika u seriji s učestalošću u standardnom turskome jeziku: točnije, uspoređeno je prvih sto najfrekventnijih lema iz istraživana korpusa sa sto najfrekventnijih lema iz referentnog izvora *A Frequency Dictionary of Turkish* (2017).

A Frequency Dictionary of Turkish autora Yeşima Aksana, Mustafe Aksana, Ümita Mersinlija i Umuta Ufuka Demirhana jedini je rječnik te vrste i jedan je od najvažnijih i najopsežnijih frekvencijskih izvora za poučavanje i istraživanje turskoga jezika. Rječnik donosi 5.000 najčešće korištenih riječi u suvremenom pisanom i govornome turskom jeziku. Učestalost se temelji na podacima iz ranije spomenutoga *Turskog nacionalnog korpusa* koji obuhvaća 50.997.016 aktivnih riječi i predstavlja širok raspon tekstualnih kategorija u razdoblju od 23 godine (1990. – 2013.). Korpus je sastavljen od pisanih podataka (98 %) te transkribiranoga govora (2 %), što omogućuje pouzdanu i reprezentativnu sliku najčešće upotrebljavanog leksika (Aksan i sur., 2017, str. 1–8).

Usporedbom popisa sto najfrekventnijih lema iz analiziranoga korpusa s popisom sto najfrekventnijih lema iz izvora *A Frequency Dictionary of Turkish* utvrđeno je podudaranje 50 lema, odnosno 50 % ukupnoga broja uspoređenih jedinica.⁷

Najveći stupanj podudarnosti zabilježen je među glagolima, što je osobito relevantno za ovu analizu. Među 31 glagolom koji se pojavljuje u popisu najfrekventnijih lema analiziranoga korpusa njih 23 nalazi se i među najfrekventnijim glagolima suvremenoga turskog jezika. Time udio podudarnih glagola iznosi približno 74,19 %, što upućuje na to da se velik dio glagola prisutnih u analiziranim epizodama serije ujedno ubraja i među najfrekventnije glagole suvremenoga turskog jezika. Glagoli koji se ponavljaju su: *olmak* 'biti', *etmek* 'činiti', *yapmak* 'raditi', *almak* 'uzeti', *demek* 'reći', *gelmek* 'doći', *vermek* 'dati', *görmek* 'vidjeti', *çıkmaq* 'izaći', *bulmak* 'pronaći', *aramak* 'tražiti', *gitmek* 'otići', *istemek* 'željeti', *geçmek* 'proći', *bilmek* 'znati', *anlamak* 'razumjeti', *kalmak* 'ostati', *söylemek* 'govoriti', *bakmak* 'gledati', *yanmak* 'gorjeti', *girmek* 'ući', *getirmek* 'donijeti', *açmak* 'otvoriti'.

Uz glagole kao najfrekventniju kategoriju u oba se korpusa ponavljaju neke osnovne imenice, primjerice *yer* 'mjesto', *zaman* 'vrijeme', *şey* 'stvar', *gün* 'dan', *el* 'ruka', *iş* 'posao', *çocuk* 'dijete', *yüz* 'lice' i *ev* 'kuća'. U korpusu serije posebno se ističu imenice povezane s obiteljskim odnosima – *abi* 'stariji brat', *abla* 'starija sestra', *anne* 'majka', *baba* 'otac', *kardeş* 'brat/sestra' i *oğul* 'sin' – koje se ne pojavljuju među prvih sto lema u *A Frequency Dictionary of Turkish*. Njihova se učestalost može povezati sa sadržajem serije.

U obje se liste frekventnih riječi ponavljaju i česte zamjenice – *bu* 'ovo', *o* 'on/ona', *ben* 'ja', *sen* 'ti', *biz* 'mi', *kendi* 'vlastiti' – zatim veznik *ama* 'ali', veznik i prilog *yani* 'dakle', prilozi *çok* 'puno', *daha* 'još', *hiç* 'ništa', 'nimalo', *böyle* 'tako', veznik i čestica *ki* 'da', 'ta', 'pa' i *ya* 'ili', 'ta', 'pa' te pridjev *iyi* 'dobar'. Ta skupina leksema predstavlja jezgreni i očekivani vokabular svakodnevnoga jezika pa se njihova visoka frekvencija u oba izvora može smatrati uobičajenom.

Na temelju provedene usporedbe može se zaključiti da je leksik analiziranih epizoda serije u velikoj mjeri podudaran s najfrekventnijim leksikom suvremenoga turskog jezika, osobito kada je riječ o glagolima i jezgrenome svakodnevnom vokabularu. Iako se zaključci odnose na ograničen frekvencijski uzorak, dobiveni rezultati sugeriraju da se jezik analizirane serije, unatoč određenim tematskim specifičnostima, temelji na čestom i komunikacijski relevantnom vokabularu suvremenoga turskog jezika.

7 Tablice s prvih sto lema u korpusu serije i sto najfrekventnijih lema iz izvora *A Frequency Dictionary of Turkish* dostupne su u Podvezanec (2025, str. 21–23). Analiza lema u tome je istraživanju provedena na poduzorku koji obuhvaća prve tri epizode serije (22.256 pojavnica), pri čemu je, zbog ograničene pouzdanosti automat-ske lematizacije za turski jezik, veći dio lematizacije proveden ručno. Iako se cjelokupno istraživanje u ovome radu temelji na korpusu prvih pet epizoda, dobiveni se rezultati mogu smatrati reprezentativnima s obzirom na to da se analizirane epizode odvijaju u istom komunikacijskom i narativnom okviru te uključuju isti skup likova i slične tipove interakcije, čime je varijabilnost leksičke distribucije razmjerno ograničena.

5. Analiza glagolskih oblika u analiziranom korpusu

5.1. Glagol imek 'biti'

U turskom jeziku funkciju i značenja pomoćnoga glagola *biti* nose glagoli *imek* i *olmak*. Glagol *imek* obuhvaća i oblike koji s njime nisu u genetskoj, nego samo u funkcionalnoj vezi. To su lični nastavci za prezent, imenska negacija *değil* te predikativi *var* i *yok* (Čaušević, 1996, str. 191). I oni su uključeni u donju analizu.

Za razliku od *olmak* koji je punoznačni odnosno leksički glagol, *imek* je nepotpun odnosno suponentan glagol, s posebnim odričnim oblicima; koristi se kao kopula, ne može se upotrijebiti kao leksički infinitiv i nije produktivan u suvremenome turskom jeziku (Čaušević i Kerovec, 2021, str. 44). Zbog svega navedenoga oblici glagola *imek* bit će obrađeni prvi, a glagol *olmak* bit će uključen u analizu s drugim punoznačnim glagolima.

5.1.1. Predikativi *var* i *yok*

Predikativima *var* i *yok* izražava se egzistiranje i posjedovanje. Na hrvatski jezik predikativ *var* prevodi se glagolima 'nalaziti se', 'postojati', 'biti', 'imati', a predikativ *yok* kao njegova negacija ('ne nalaziti se' itd.).

Tablica 1. Morfološki oblici predikativa *var* i *yok*⁸

oblik	var	broj pojava	yok	broj pojava
prezent 1JD ⁹	varım	3	yokum	1
prezent 2JD	varsın	1	/	/
prezent 3JD	var	181	yok	241
prezent 1MN ¹⁰	varız	3	/	/
prezent 2MN	varsınız	1	/	/
perfekt 2MN	vardın	1	/	/
perfekt 3JD	vardı	20	yoktu	6
modalnost – <i>Dir</i> 3JD	vardır	8	yoktur	1
modalnost – <i>mİş</i> 1JD	/	/	yokmuşum	1
modalnost – <i>mİş</i> 3JD	varmış	6	/	/
kondicional 3JD	varsa	5	yoksa	15
konverb <i>i-ken</i>	varken	1	yokken	1

⁸ U tablici nisu prikazane potpune paradigme, nego samo oni oblici koji su potvrđeni u analiziranom korpusu.

⁹ Kratica za: 1. lice jednine.

¹⁰ Kratica za: 1. lice množine.

Oba predikativa najbrojniji su u prezentu 3. l. jednine (*var/yok*). Unutar toga broja nalazi se i 19 slučajeva kad se predikativ *var* pojavljuje u upitnom obliku (*var mı*) te 15 slučajeva u kojima se *yok* pojavljuje u upitnom obliku (*yok mu*). Od ostalih oblika u značajnijem broju pojavljuju se samo kondicional 3. l. jednine (*varsas/yoksa*) te perfekt 3. l. jednine (*vardı/yoktu*).

5.1.2. Prezent glagola *imek*

Kako je rečeno, prezent glagola *imek* tvori se ličnim nastavcima koji s njim nisu u genetskoj vezi. U korpusu se pojavljuje u sljedećem broju primjera:¹¹

Tablica 2. Afirmativne paradigme prezenta glagola *imek*

	afirmativna paradigma	broj pojava	afirmativno-upitna paradigma	broj pojava
jednina	1. –(y)Im	68	1. mIyIm	7
	2. –sIn	38	2. mIsIn	45
	3. –DIr / –ø	61 x ¹² –DIr u cijelom korpusu, te 19 x –ø na 3 kartice teksta ¹³	3. mIDIr / mI	mIDIr x 2, mI x 256
množina	1. –yIz	29	1. mIyIz	9
	2. –sInIz	10	2. mIsInIz	6
	3. DIr–lAr / –DIr / –lAr	3	3. mIDIrlAr / mIDIr / –lAr mI	0

U usvajanju ovih sufikasa i oblika vjerojatno mogu pomoći ustaljeni izrazi svakodnevnih komunikacija koji se često ponavljaju, kao što su npr. *iyiyim* ‘dobro sam’ (16x) i *üzgünüm* ‘žao mi je’ (5x), *iyi misin* ‘jesi li dobro?’ (21x), *öyle mi* ‘tako?’ (10x), *tamam mı* ‘u redu?’ (43x) te fraza *hayırdır* ‘je li sve u redu?’ (19x).

11 Budući da za turski jezik u sustavu *Sketch Engine* nije dostupna potpuna morfološka analiza ni lematizacija, oblici pomoćnog glagola *imek* morali su biti pretraživani ručno. U nedostatku mogućnosti automatskog grupiranja prema lemi svi su oblici toga glagola traženi pojedinačno, prema svim njihovim fonološkim varijantama i alomorfima (npr. samo za 3JD to su sufiksi: –*dim*, –*dim*, –*dum*, –*düm*, –*tım*, –*tım*, –*tum*, –*tüm*). Ujedno je bilo potrebno ručno provjeriti svaki kontekst kako bi se isključili oblici koji pripadaju drugim glagolima, a ne glagolu *imek*.

12 Oznaka „x“ označuje broj ponavljanja („puta“).

13 Slučajevi u kojima je kopula –*DIr* ispuštena veoma su česti, a kako nisu pretraživi preko *Sketch Enginea* (jer sufiksa nema), pobrojan je broj pojava na prve tri kartice teksta korpusa kao ilustracija. Kako cijeli korpus ima 97 kartica, može se pretpostaviti da u cijelom korpusu ima približno 614 ovakvih slučajeva.

Tablica 3. Negativne paradigme prezenta glagola *imek*

	negativna paradigma	broj pojava	negativno–upitna paradigma	broj pojava
jednina	1. değilim	18	1. değil miyim	0
	2. değilsin	6	2. değil misin	2
	3. değil / değildir	değil x 93, değildir x 0	3. değil mi	21
množina	1. değiliz	9	1. değil miyiz	0
	2. değilsiniz	3	2. değil misiniz	0
	3. değiller	2	3. değiller mi	0

Broj primjera iz negativnih paradigmi prezenta glagola *imek* znatno je manji od afirmativnih. Struktura *değil mi?* može značiti i upitni izraz ‘zar ne?’, i u tom se značenju u korpusu pojavljuje 68 puta.

5.1.3. Perfekt glagola *imek*

Primjeri iz negativnih paradigmi pojavljuju se ukupno u devet slučajeva: 3JD *değildi* x 8 i 1JD *değildim* x 1. Negativno–upitni oblici ne pojavljuju se.

Tablica 4. Afirmativne paradigme perfekta glagola *imek*

	afirmativna paradigma	broj pojava	afirmativno–upitna paradigma	broj pojava
jednina	1. –Dİm	8	1. mİydİm	0
	2. –Dİn	1	2. mİydİn	2
	3. –Dİ	51	3. mİydİ	4
množina	1. –Dİk	0	1. mİydİk	0
	2. –Dİnİz	5	2. mİydİnİz	2
	3. –DİAr	1	3. mİydİAr	0

5.1.4. Evidencijal glagola *imek*

U odnosu na prethodna ovo se glagolsko vrijeme pojavljuje u manjem broju slučajeva. Od afirmativno–upitnih oblika pojavljuje se samo 3JD *mİymİş* x 3, a od negativnih paradigmi samo 3JD *değilmiş* x 1.

Tablica 5. Afirmativna paradigma evidencijala glagola *imek*

	afirmativna paradigma	broj pojava
jednina	1. –mİşIm	1
	2. –mİşsIn	3
	3. –mİş	36
mnöžina	1. –mİşIz	1
	2. –mİşsInIz	0
	3. –mİşlAr	0

5.1.5. Konverb *i–ken* i kondicional glagola *imek*

Od ukupno 39 primjera konverba *i–ken* u cijelom korpusu 11 primjera tvoreno je preko glagola *imek*. Kondicional glagola *imek* znatno je manje zastupljen. Npr. u cijelom korpusu za 1JD prisutna su 23 primjera kondicionala i kondicionalne modalnosti, i nijedan od tih primjera nije tvoren preko glagola *imek*; za 2JD od 39 primjera samo su 2 glagola *imek*.

5.2. Ostali glagoli

Zbog ograničenih mogućnosti strojnog pretraživanja morfoloških obilježja u analizi punoznačnih glagola nisu uzeti svi glagoli, nego deset najfrekventnijih, navedenih u tablici 6.

Tablica 6. Analizirani glagoli s brojem pojava i različenica, složeni prema broju pojava

glagol	broj pojava	broj različenica
<i>olmak</i> ‘biti’, ‘postati’	775	101
<i>gelmek</i> ‘doći’	418	72
<i>yapmak</i> ‘(u)činiti’	379	94
<i>etmek</i> ‘činiti’	322	83
<i>bakmak</i> ‘gledati’	307	36
<i>gitmek</i> ‘ići’	301	59
<i>almak</i> ‘uzeti’	221	60
<i>bilmek</i> ‘znati’	174	44
<i>vermek</i> ‘dati’	164	56
<i>çıkma</i> ‘izaći’	150	49

Glagoli *olmak* ‘biti’, ‘postati’ i *etmek* ‘činiti’ pomoćni su glagoli pa je njihova učestalost (i pojava i različica) i očekivana. Brojnost glagola *gelmek* ‘doći’, *yapmak* ‘(u)činiti’ i *gitmek* ‘ići’ može se objasniti njihovom važnošću za organizaciju i odvijanje radnje. Kod glagola *bakmak* ‘gledati’ uočljiva je veća razlika u broju pojava i različenica, uzrokovana time što se 177 puta pojavljuje kao izraz *bak* u funkciji pozivnog ili fokusnog markera ‘vidi’, ‘slušaj’, ‘gle’, i često se koristi u govornom jeziku. U odnosu na glagole *vermek* ‘dati’ i *çikmak* ‘izaći’ s kojima ima sličan broj pojava, glagol *bilmek* ‘znati’ ima nešto manji broj različenica jer se 33 puta pojavljuje u obliku *biliyor* ‘zna’ i 24 puta kao *bilmiyorum* ‘ne znam’. Uglavnom, manji udio različenica u odnosu na broj pojava upućuje na to da se glagol pojavljuje u ograničenom broju standardiziranih i često ponavljanih oblika, dok veći udio znači veću morfološku raznolikost.

U analizi koja slijedi obuhvaćene su one kategorije glagolske morfologije (lice i broj, vrijeme, negacija i sl.) koje su smatrane relevantnima. Primjerice, kod jednostavnih vremena analiziran je udio prema paradigmama jer se na tom velikom dijelu korpusa htjelo pokazati kakva je raspodjela primjera, a kako je kod složenih glagolskih vremena distribucija tih kategorija vrlo slična, nije ponovno prikazivana.

5.2.1. Jednostavna glagolska vremena

U jednostavna glagolska vremena ubrajaju se prezent na *–(I)yor*, prezent na *–(*)r*, futur, perfekt na *–mİş* i perfekt na *–DI* (Čaušević, 1996, str. 236–267). Raspodjela njihove zastupljenosti u korpusu jest sljedeća:

Tablica 7. Raspodjela zastupljenosti jednostavnih glagolskih vremena¹⁴

glagolsko vrijeme	broj	udio
perfekt na <i>–DI</i>	484	36,54 %
prezent na <i>–(I)yor</i>	472	35,65 %
prezent na <i>–(*)r</i>	277	20,94 %
futur	197	14,88 %
perfekt na <i>–mİş</i>	91	6,87 %

Visoki udio perfekta na *–DI* kojim se izražava prošla svršena i posvjedočena radnja i prezenta na *–(I)yor* kojim se izražava tekuća, „prava“ sadašnjost očekivan je s obzirom na radnju i događaje u kontekstu serije. Prezent na *–(*)r* svezvremenski je i futurski prezent, ali se njime izriču i molbe, zahvale i želje, što znatno utječe na broj njegovih pojava. Futur i perfekt na *–mİş*, kao modalno i pragmatički obilježena vremena, očekivano su manje zastupljeni: futur jer je radnja u seriji usmjerena ponaj-

14 Za tehničku pomoć u izradi tablica i računanju postotaka korišten je jezični model umjetne inteligencije (ChatGPT).

prije na sadašnjost i neposrednu prošlost, a perfekt na *-mİş* jer svojim modalnim značenjima (neosvjedočena radnja, zaključivanje, pretpostavka) nije prikladan za prikaz događaja koji su u seriji jasno prikazani i izravno praćeni.

Brojčana raspodjela zastupljenosti primjera po paradigmama i po vremenima može se vidjeti u sljedećoj tablici:

Tablica 8. Zastupljenost primjera glagolskih vremena prema paradigmama

glagolsko vrijeme	afirmativna paradigma	afirmativno–upitna paradigma	negativna paradigma	negativno–upitna paradigma	ukupno
prezent <i>–(I)yor</i>	360	24	85	3	472
prezent <i>–(*)r</i>	184	38	54	1	277
futur <i>–AcAK</i>	172	4	19	2	197
perfekt <i>–mİş</i>	73	2	16	0	91
perfekt <i>–DI</i>	440	15	27	2	484
ukupno	1229	83	201	8	1521

Najbrojniji su primjeri afirmativne paradigme i nakon toga negativne, dok je znatno manji broj primjera upitnih paradigmi. Veći broj upitnih oblika u prezentu na *–(*)r* može se objasniti time što se njime upitnim rečenicama izražava molba.

Raspodjela primjera prema licima jest pak sljedeća:

Tablica 9. Raspodjela primjera prema licima

lice i broj	ukupno	udio (%)
3JD	727	47,8
1JD	332	21,8
2JD	259	17,0
1MN	106	7,0
2MN	53	3,5
3MN	44	2,9

Primjeri u 3. licu jednine najzastupljeniji su zbog naracije i opisivanja radnje, primjeri 1. i 2. lica jednine brojni su zbog dijaloške strukture izgovorenoga teksta, dok su oblici u množini očekivano manje brojni.

5.2.2. Složena glagolska vremena

Složena glagolska vremena, imperfekt na $-(I)yordu$ i na $-rdI$, pluskvamperfekt na $-mIş-tI$ i na $-DI-yDI$ te futur perfekta na $-(y)AcAk-tI$ (Čaušević, 1996, str. 267–286), zastupljena su u znatno manjem broju u odnosu na jednostavna.

Tablica 10. Zastupljenost složenih glagolskih vremena i raspodjela prema licima

lice i broj	imperfekt na $-(I)yordu$	imperfekt na $-rdI$	pluskvamperfekt na $-mIş-tI$	pluskvamperfekt na $-DI-yDI$	futur perfekta na $-(y)AcAk-tI$	ukupno
1JD	9	6	3	/	3	21
2JD	4	/	1	/	3	8
3JD	8	5	6	/	4	23
1MN	2	/	2	/	/	4
2MN	/	/	/	/	/	0
3MN	/	/	2	/	/	2
ukupno	23	11	14	0	10	58

Kako se vidi, najbrojnije složeno vrijeme jest imperfekt na $-(I)yordu$, a raspodjela broja primjera prema licima podudarna je s onom kod jednostavnih vremena.

5.2.3. Glagolski načini

Ovdje će biti obrađeni imperativ, optativ, necesitativ i kondicional¹⁵ (Čaušević, 1996, str. 288–302), kod kojih se modalnost izražava posebnim sufiksima.

Tablica 11. Broj pojava imperativa, optativa, necesitativa i kondicionala

glagolski način	afirmativni oblici	negativni oblici	ukupno
imperativ	723	120	843
optativ	147	4	151
necesitativ	1	/	1
kondicional	40	9	49

Među analiziranim glagolskim načinima daleko je najbrojniji imperativ. Kao relevantan podatak navedeni su i udjeli primjera negativnih tvorbi jer se može pret-

¹⁵ Pod kondicionalom u ovom slučaju podrazumijevamo i primjere kondicionala i primjere kondicionalne modalnosti koja se tvori dodavanjem kondicionalnog sufiksa na participske i modalne osnove.

postaviti da znatno utječu na usvajanje korištenja negacijskog sufiksa *-mA* kojim se negiraju svi finitni i infinitni glagolski oblici (osim glagola *imek*).

Tablica 12. Raspodjela primjera korištenja glagolskih načina prema licima

lice i broj	imperativ	kondicional	optativ	zbroj
1JD	0	10	61	71
2JD	664	28	0	692
3JD	90	5	0	95
1MN	0	0	90	90
2MN	84	3	0	87
3MN	5	3	0	8

Visoka frekvencija oblika u 2. licu upućuje na dijalošku strukturu scenarija u kojemu je izravno obraćanje među likovima jedan od ključnih nositelja radnje. Istodobno, znatan udio oblika u 1. licu (i jednine i množine) odražava subjektivizaciju. Oblici u 3. licu u ovome slučaju ne ukazuju na narativne dijelove, nego – s obzirom na to da je riječ o imperativnoj paradigmi – funkcioniraju kao naredbe i poticaji upućeni trećoj osobi, čime izravno doprinose pokretanju i dinamiziranju radnje.

5.2.4. Ostali načini izražavanja modalnosti¹⁶

U korpusu analiziranih deset glagola zastupljeno je 5 primjera modalnosti na *-DIr* i 7 primjera modalnosti na *-mİş*.

Od analitičkih formi izrazito su brojni posibilitativ (*-(y)A-bilmek*) i imposibilitativ (*-(y)A-mA-mAk*). Posibilitativ se pojavljuje u 48 slučajeva, a imposibilitativ u 52 primjera. Može se, dakle, reći da se češće izražava nemogućnost nego mogućnost radnje, tim više što se u navedenom broju korištenja posibilitativa 20 slučajeva odnosi na izraz *olabilir* (posibilitativ glagola *olmak* u prezentu na *-(*)r*) u značenju hrvatskog priloga 'moguće'). Od analitičkih formi tvorenih drugim okazionalno-pomoćnim glagolima u korpusu su zastupljena dva primjera, oba s glagolom *vermek* 'dati', u značenju brzine i trenutačnosti glagolske radnje (npr. *gidiver* 'skokni', 'odju-ri').

Za analizu perifrastičnih formi analiza deset glagola nije bila dovoljna, nego se pretraga morala proširiti na cijeli korpus serije. Nije bilo moguće sve pretražiti ručno, ali su npr. već pregledom konkordancija prvih triju lica perfekta na *-DI* glagola *olmak* pronađena tri primjera perifrastičnih formi (*anlamış oldum* 'razumio sam', *duymuş oldun* 'čuo si', *bıçaklamış oldun* 'izbo si nožem', što ukazuje na to da bi ih moglo biti i s drugim oblicima glagola *olmak*. Primjeri izražavanja prividnoga vršenja glagolske radnje koji se tvore uvođenjem postpozicije *gibi* također su prisut-

16 O ostalim načinima izražavanja modalnosti vidjeti u: Čaušević (1996, str. 300–319).

ni: u cijelom analiziranom korpusu postpozicija *gibi* pojavljuje se 130 puta, od toga čak 17 puta unutar perifrastičnih formi (npr. *ben burada yokmuşum gibi* 'kao da me nema ovdje' i sl.).

5.2.5. Participi¹⁷

Najbrojniji među participima je particip na $-(y)An$, a među primjerima toga participa najčešći su oni tvoreni od glagola *olmak* 'biti'. Među tim primjerima glagola *olmak* zabilježena su i dva složena participa (*verecek olan* 'koji će/bi mogao/bi trebao dati' i *solmuş olan* 'koji je uvenuo').

Tablica 13. Zastupljenost primjera participa u korpusu

particip	broj primjera
$-(y)An$	53
$-(*)r$	0
$-(y)AcAk$	2
$-mIş$	3

5.2.6. Glagolske imenice na $-mAk$ i $-mA$ ¹⁸

U korpusu analiziranih glagola zastupljeno je 18 primjera glagolske imenice $-mAK$. Sve su u apsolutnom padežu i uglavnom su dio imenskog predikata *gerek/lâzım olmak* ili u genitivnoj vezi s lokativom imenice *zor* 'prisila', u značenju potrebe (npr. *dikkat etmek lazım* 'potrebno je paziti', *kabul etmek zorundayız* 'moramo prihvatiti' i sl.). Nadalje, glagolska imenica na $-mAK$ u tri se slučaja koristi s postpozicijom *için* tvoreći adverbijal namjere te dva puta s prilogom *birlikte*, u značenju adverbijala koncesivnosti.

Glagolska imenica na $-mA$ bez dodanih posvojnih sufikasa u dva slučaja koristi se u apsolutnom padežu kao dio genitivne veze (npr. *çıkma ihtimali var* 'postoji mogućnost pojavljivanja'). U kosim padežima glagolska imenica na $-mA$ bez dodanih posvojnih sufikasa dva se puta pojavljuje u akuzativu, u funkciji objekta (npr. *bırak oğlum olmayı, kapımdaki köpeğim olamazsın* 'ne možeš biti ni pas pred mojim vratima, a kamoli moj sin'), a u preostalih 17 slučajeva imenica je u dativu, u značenju adverbijala namjere (npr. *hava almaya çıktım* 'izašao sam uzeti zraka') ili s faznim glagolima (npr. *bazı şeyler gelmeye başladı gözümün önüne* 'pred očima su mi se počele pojavljivati neke stvari'). U ostalih 17 primjera glagolska imenica na $-mA$ ima i dodane posvojne sufikse i ima uobičajene rečenične funkcije. U analiziranom korpusu nema primjera u kojima se glagolska imenica na $-mA$ pojavljuje s postpozicijama, kao semantički ekvivalent hrvatskih zavisnih priložnih rečenica.

17 O participima vidjeti u: Čaušević (1996, str. 319–328).

18 O glagolskim imenicama na $-mAk$ i $-mA$ vidjeti u: Čaušević (1996, str. 328–344).

5.2.7. Glagolske imenice odnosno proparticipi na *-DIk* i *-(y)AcAk*¹⁹

Raspodjela navedenih glagolskih oblika u analiziranom korpusu glagola jest sljedeća:

Tablica 14. Zastupljenost primjera proparticipa u korpusu

	<i>-DIk</i>	<i>-(y)AcAk</i>	ukupno
imenička funkcija	52	15	67
atributna funkcija	33	6	39
ukupno	85	21	106

Prema dobivenim podacima može se vidjeti da se glagolske imenice na *-DIk* i *-(y)AcAk* gotovo dvostruko više pojavljuju u imeničkoj nego u atributnoj funkciji. Također, vidi se da je broj primjera glagolske imenice na *-DIk* četiri puta veći od one na *-(y)AcAk*.

5.2.8. Konverbi i kvazikonverbi²⁰

Kako se vidi iz tablice 15, najbrojniji je konverb *-(y)Ip*. Nakon njega slijedi konverb *-(y)ArAk*, međutim, od njegova 21 primjera, 19 se odnosi na *olarak*, konverbni oblik glagola *olmak* 'biti' koji znači 'kao' (dosl. 'bivajući kao', 'poput') i kojim se uvode priložne oznake načina i popratnih okolnosti (Čaušević i Kerovec, 2021, str. 196). Konverb *-(y)All* se pak svih pet puta pojavljuje u izrazu *bildim bileli* 'oduvijek', 'otkako znam'.²¹

Tablica 15. Zastupljenost primjera konverba i kvazikonverba

konverb	broj primjera	kvazikonverb	broj primjera
<i>-(y)Ip</i>	40	<i>-mA+dAn</i> ²²	3
<i>-(y)ArAk</i>	21	<i>-(*)r-mAz</i>	2
<i>-(y)IncA</i>	10	<i>-CasInA</i>	0
<i>i-ken</i>	5	<i>-DIkçA</i>	0
<i>-mAdAn</i>	5	<i>-DIktA</i>	0
<i>-(y)All</i>	5	<i>-DIktAn başka</i>	0
		<i>-DIktAn sonra</i>	0

19 O glagolskim imenicama (proparticipima) na *-DIk* i *-(y)AcAk* vidjeti u: Čaušević (1996, str. 344–371).

20 O konverbima i kvazikonverbima vidjeti u: Čaušević (1996, str. 372–399).

21 <https://lugatim.com/s/bilmek>, 24.10.2025.

22 S postpozicijom *önce*.

5.2.9. Zastupljenost glagolskih oblika u korpusu

Nakon što su detektirani i prebrojeni svi primjeri glagolskih oblika, moguće je izraditi pregled njihove relativne zastupljenosti prema učestalosti. U prikazani su poredak uključeni samo oni glagolski oblici koji u korpusu imaju više od deset zabilježenih primjera. Potrebno je ponovno naglasiti da je analiza punoznačnih glagola provedena samo na deset najfrekventnijih glagola u korpusu te da bi se u cijelom korpusu apsolutni broj njihovih kategorija zasigurno pokazao većim. Unatoč tomu, posebno je zanimljiva vrlo visoka učestalost glagola *imek* 'biti' te pripadajućih predikativnih izraza *var* i *yok*. Potrebno je također svratiti pažnju na to da udio imperativa dvostruko nadilazi udjele glagolskih vremena. Infinitne se forme, očekivano, nalaze među manje brojnim kategorijama jer se tekst uglavnom gradi na finitnim formama, odnosno jednostavnim rečenicama.

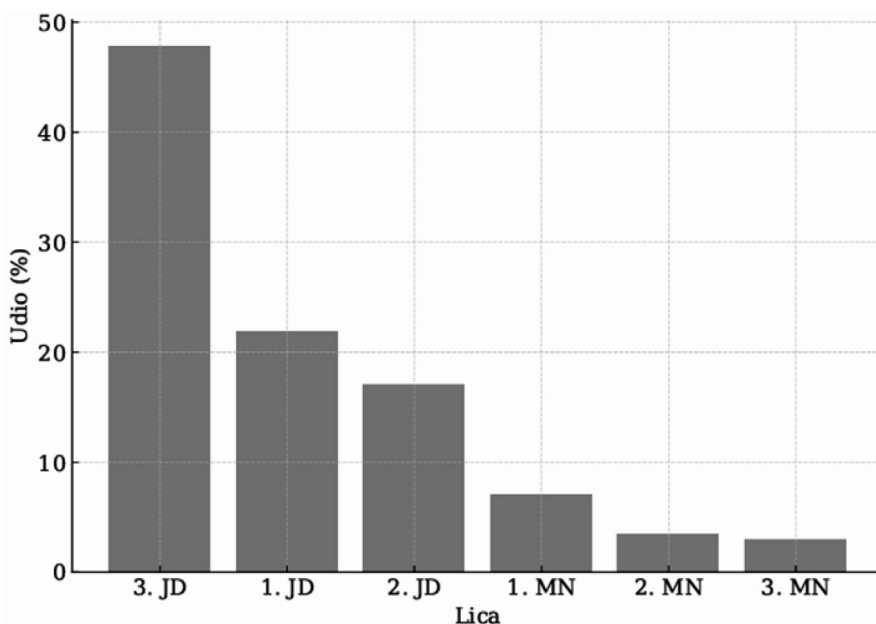
Tablica 16. Poredak glagolskih oblika prema broju pojava

redni broj	broj pojava	glagolska kategorija
1	1.302	<i>imek</i> – prezent
2	843	imperativ
3	484	perfekt na <i>-DI</i>
4	472	prezent na <i>-(I)yor</i>
5	277	prezent na <i>-(*)r</i>
6	241	<i>yok</i> – prezent
7	197	futur
8	181	<i>var</i> – prezent
9	151	optativ
10	91	perfekt na <i>-mIş</i>
11	85	glagolska imenica <i>-DIk</i>
12	83	<i>imek</i> – perfekt
13	53	particip <i>-(y)An</i>
14	52	imposibilitativ
15	49	kondicional
16	48	posibilitativ
17	45	<i>imek</i> – evidencijal
18	40	konverb <i>-(y)Ip</i>
19	38	glagolska imenica <i>-mA</i>
20	23	imperfekt na <i>-(I)yordu</i>
21	23	glagolska imenica <i>-mAK</i>
22	21	glagolska imenica <i>-(y)AcAk</i>

23	21	konverb $-(y)ArAk$
24	20	<i>var</i> – perfekt
25	17	perifrastične forme + <i>gibi</i>
26	15	<i>var</i> – kondicional
27	14	pluskvamperfekt na $-mIş-tI$
28	11	<i>imek</i> – <i>i-ken</i>
29	10	futur perfekta na $-(y)AcAk-tI$
30	10	konverb $-(y)IncA$

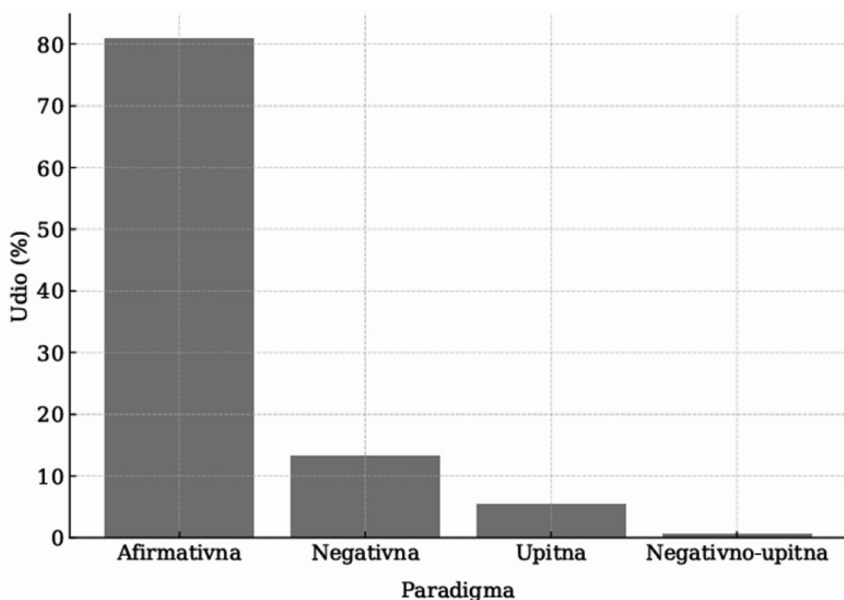
Kako je već napomenuto, zastupljenost glagolskih lica analizirana je samo u sklopu jednostavnih glagolskih vremena. Kod ostalih kategorija raspodjela je vrlo slična (izuzev imperativa i optativa, v. 5.2.3.), pa se može smatrati da je ova analiza reprezentativna za cijeli korpus. Iz odnosa među licima vidljivo je da je 3. lice jednine najbrojnije, slijede ga 1. i 2. lice jednine, dok su oblici množine znatno rjeđi.

Slika 1. Zastupljenost glagolskih lica u jednostavnim vremenima (u %)



Jednako tako, na temelju analize jednostavnih glagolskih vremena može se zaključiti da je među paradigmama daleko najzastupljenija afirmativna paradigma, dok su ostale – upitna, negativna i negativno-upitna – zastupljene u znatno manjim omjerima.

Slika 2. Udio glagolskih paradigmi u analiziranome korpusu (u %)



6. Gledanje serija kao izvor *inputa*: istraživanje i analiza

Nakon provedene korpusne analize glagolskih oblika u jeziku televizijske serije, u ovome se dijelu rada razmatra i pitanje njihove moguće prepoznatljivosti među gledateljima izloženima turskom jeziku putem serija. Polazište ovoga dijela istraživanja oblikovano je prema referentnim razinama poznavanja jezika opisanima u Zajedničkom europskom referentnom okviru za jezike (ZEROJ). Budući da za turski jezik ne postoji određeni referentni dokument koji bi definirao sadržaje pojedinih razina, svaki udžbenički niz u skladu sa ZEROJ-em razvija vlastiti sadržaj u okviru odgovarajuće referentne razine poučavanja i učenja jezika (v. Erdil i Açıık, 2021). Kao temelj za sastavljanje upitnika u ovome istraživanju odabran je udžbenički niz *Yeni Hitit* (Gazi Üniversitesi TÖMER, 2016a, 2016b) jer je riječ o nastavnom materijalu koji je razvila škola turskoga jezika Ankara TÖMER koja je institucija s dugom tradicijom i visokom razinom prepoznatljivosti u području poučavanja turskoga kao stranoga jezika.

U istraživanju je sudjelovalo 15 ispitanika: 7 studenata prve godine studija Turkologije na Filozofskom fakultetu Sveučilišta u Zagrebu te 8 polaznika tečaja turskog jezika razine A1 pri Institutu Yunus Emre u Zagrebu.²³ Uzorak je formiran na temelju dostupnosti ispitanika u navedenim institucijama koje su, koliko nam je poznato, jedina mjesta organiziranoga poučavanja turskoga jezika u Republici Hr-

23 Zahvaljujemo Institutu Yunus Emre Institute u Zagrebu na omogućavanju provedbe istraživanja te posebno voditeljici tečaja, magistri turkologije Andrei Reljac.

vatskoj. Istraživanje je provedeno početkom studenoga 2025. godine, u razdoblju kada su ispitanici tek započeli formalno učenje turskoga jezika. Uključeni su samo oni studenti i polaznici koji su sami procijenili da su gledanjem turskih televizijskih serija stekli određenu razinu poznavanja jezika. Istraživanje je bilo kvantitativno i sastojalo se isključivo od pismenoga dijela: upitnik je sadržavao zadatke višestrukog izbora s jednim točnim odgovorom, bez uključivanja kvalitativnih pitanja. Iako je uzorak malen, dobiveni nalazi mogu se smatrati indikativnima te poslužiti kao smjernica za daljnja, opsežnija istraživanja.

Prije pitanja za provjeru jezičnih znanja ispitanici su odgovorili na nekoliko uvodnih pitanja. Prema njima se vidi da je samo jedan ispitanik prethodno pohađao tromjesečni *online* tečaj turskog jezika na razini A1, dok su svi ostali turski usvajali isključivo gledanjem serija. Četiri ispitanika turske serije gledaju u razdoblju dugom od 4 do 6 godina, a njih deset više od 6 godina. Nadalje, 5 ispitanika serije gleda jednom tjedno, dok ih 9 gleda više puta tjedno. Što se tiče načina gledanja, serije gledaju s hrvatskim titlovima (11 ispitanika), s engleskim titlovima (3), bez titlova (8) te s turskim titlovima (3).²⁴ U procjeni vlastitoga znanja turskog jezika 5 ispitanika navelo je da pasivno razumije pojedine riječi, 4 smatra da razumije osnovne rečenice i fraze, 1 ispitanik dobro razumije dijelove dijaloga, 5 ispitanika vrlo dobro razumije većinu dijaloga, a 2 ispitanika izjavila su da razumiju gotovo sve. Kao glavne motive za gledanje turskih serija ispitanici su naveli jezik (11 ispitanika), kulturu i tradiciju (11), glumce i radnju (9) te glazbu i ambijent (7).

U provjeri jezičnih znanja upitnik se sastojao od 28 pitanja raspoređenih prema razinama A1 – B2,²⁵ po 7 pitanja za svaku razinu.²⁶ Svako se pitanje sastojalo od rečenice na hrvatskom jeziku, a ispitanici su morali među četiri ponuđena odgovora prijevoda na turski odabrati po jedan odgovor.

24 U ovom i sljedećim pitanjima ispitanici su mogli izabrati više opcija. U istraživanjima audiovizualnoga inputa naglašava se da titlovi na jeziku koji se uči mogu olakšati povezivanje govornoga i pisanoga oblika jezika te pridonijeti lakšemu uočavanju nepoznatih riječi i jezičnih obrazaca, dok titlovi na materinskome jeziku, premda ne usmjeravaju pozornost jednako snažno na jezični oblik, i dalje predstavljaju važan izvor jezičnoga inputa. Istodobno se ističe i važnost vizualnoga konteksta, odnosno slike i radnje koje pružaju dodatne tragove za razumijevanje jezičnih elemenata u audiovizualnome sadržaju (Montero Perez, Sydorenko i Huntley, 2024, str. 40–41).

25 Razine C1 i C2 nisu bile obuhvaćene provjerom jer bi izrada pitanja na tim razinama zahtijevala znatno složenije kontekste, dulje tekstove i opširnije zadatke, što bi premašilo okvir ovoga upitnika. Osim toga, pitanja do razine B2 omogućuju vrlo jasno uočavanje razlika u usvajanju finitnih i infinitnih glagolskih formi, što je posebno relevantno u učenju i poučavanju, odnosno usvajanju turskoga jezika.

26 Iako je broj pitanja ograničen, cilj nije bio iscrpiti sve gramatičke kategorije pojedine razine, nego provjeriti reprezentativni dio struktura relevantnih za istraživačka pitanja. Sedam pitanja po razini bilo je praktičan, ali i dovoljno informativan okvir.

Tablica 17. Sadržaj pitanja iz upitnika s brojem točnih odgovora i postotkom točnih odgovora prema jezičnim razinama

razina	sadržaj pitanja	broj točnih odgo- vora	postotak točnosti jezične razine
A1	1. <i>imek</i> prezent	14	58,10 %
	2. predikativ <i>var</i>	10	
	3. prezent –(I) <i>yor</i>	11	
	4. predikativi <i>var/yok</i>	3	
	5. – <i>mAdAn önce / –(I)yordu</i>	9	
	6. optativ	8	
	7. imperativ	6	
A2	1. – <i>mAdAn</i> , imposibilitativ	7	42,86 %
	2. –(y) <i>Ip</i>	6	
	3. prezent –(*) <i>r</i>	11	
	4. perfekt – <i>mİş</i>	6	
	5. – <i>mA</i>	4	
	6. perfekt – <i>mİş</i> , posibilitativ	6	
	7. futur perfekta / perfekt – <i>mİş</i>	5	
B1	1. necesitativ	6	52,38 %
	2. – <i>Diğİ</i>	6	
	3. – <i>mA</i> + posvojnost	7	
	4. kondicional	11	
	5. –(y) <i>An</i>	5	
	6. – <i>DIğİnDA</i>	11	
	7. – <i>DIğİ halde</i>	9	
B2	1. –(y) <i>AcAğİ</i> , pluskvamper- fekt – <i>mİştİ</i>	6	40,95 %
	2. – <i>DIğİ sürece</i>	7	
	3. – <i>DIğİ</i> , kondicional	10	
	4. – <i>mAk yerine</i> , – <i>mA</i>	7	
	5. – <i>DIğİ kadar</i>	6	
	6. –(y) <i>AcAğİnA</i>	3	
	7. – <i>mAsİ durumunda</i>	4	

Kako se vidi, postoci točnih odgovora prema jezičnim razinama ne smanjuju se proporcionalno prema višim razinama. No, razina A1 ima uvjerljivo najviše točnih odgovora. Na toj se razini dva zadatka odnose na predikative *var* i *yok*. U zadatku A1–2 ispitanici su trebali izabrati točan prijevod rečenice kojom se izražava egzistiranje: „Na stolu su dvije čaše i njegova kapa“ i tu su s visokim postotkom izabrali točno rješenje b): *Masada iki bardak bir de şapkası var*.²⁷ No, u zadatku A1–4, „Ali ima majku, ali nema oca“ (*Ali'nin annesi var ama babası yok*) predikativima se izražava posjedovanje pri čemu gramatički subjekt mora biti genitivna veza (dosl. „Alijeva majka postoji, ali njegov otac ne postoji“), što je očito znatno složenija struktura za usvajanje pa je broj točnih odgovora samo 3.²⁸

S obzirom na to da je imperativ u analiziranom korpusu prema učestalosti druga glagolska kategorija, dakle vrlo česta kategorija, donekle začuđuje činjenica da je zadatak A1–7 imao malen broj točnih odgovora (6). Ispitanici su trebali izabrati točan prijevod za rečenice „Neka Ahmet otvori prozor! Ti zatvori vrata!“. 6 ih je izabralo točno rješenje b): *Ahmet pencere açsın! Sen kapıyı kapat!*, a 5 ispitanika izabralo je odgovor d): *Ahmet pencere aç! Sen kapıyı kapat!* ‘Ahmete, otvori prozor! Ti zatvori vrata!’, rješenje u kojemu su oba imperativna oblika u 2. licu jednine. U analiziranom korpusu primjeri 2. lica jednine imperativa koriste se 664 puta, a u 3. licu jednine 90 puta. Moguće da zbog takve raspodjele među licima gledatelji samo 2. lice jednine prepoznaju kao imperativ.

U analiziranom korpusu kondicional nema velik broj pojava, ali u rezultatima istraživanja ima vrlo visoku točnost: kod B1–3 10 točnih odgovora i kod B2–3 10 točnih odgovora. Moguće je da na njegovu prepoznatljivost utječe posebnost njegova sufiksa koji ne nalikuje drugim sufiksima (–sA) ili to što se koristi u rečenicama koje su često nabijene emocijama.

Kod jednostavnih glagolskih vremena prezent –(i)yor (A1–3) i prezent na –(*)r (A2–3) imaju visoku točnost (oba po 11 točnih odgovora). No, to ne vrijedi za perfekt –mİş: za rečenice koje bi trebale biti prevedene tim vremenom ispitanici više biraju perfekt na –DI. To se s jedne strane može objasniti čestotnošću: u našem korpusu perfekt na –DI pojavljuje se 484 puta, a perfekt na –mİş 91 put. Osim toga, moguće je da ispitanici iz serija naprosto nisu uspjeli shvatiti koja su značenja perfekta na –mİş (neuključenost govornika u tok radnje i dr.), što je i razumljivo s obzirom na to da ni u hrvatskom ni u drugim europskim jezicima ne postoji vrijeme koje po značenju odgovara perfektu na –mİş (Čaušević, 1996, str. 255), pa su perfekatska značenja više bili skloni izraziti perfektom na –DI.

Zadatke koji su uključivali infinitne glagolske forme ispitanici su rješavali s uspjehom znatno nižim od ukupnog prosjeka. Primjerice, u A2–2 zadatak s –(y)Ip imao je samo 6 točnih odgovora, a u A2–5 zadatak s glagolskom imenicom na –mA

27 Preostala tri rješenja bila su gramatički neispravna: a) *Masada iki bardak bir de şapkası*, c) *Masada iki bardak bir de şapkası var misin*, i d) *Masada iki bardak bir de şapkan varıyor*.

28 Ostali ponudeni odgovori bili su: a) *Ali annesi var ama babası yok*, b) *Ali anne var ama baba yok* i c) *Ali annen var ama baban yok*. Ispitanici su birali i odgovore a) i b), premda su oni gramatički neispravni. Odgovor c) gramatički je ispravan, ali ne predstavlja prijevod zadane rečenice jer znači „Ali, imaš majku, ali nemaš oca“.

svega 4. Na razini B1 zadatak s proparticipom *-DIğI* (B1–2) točno je riješilo 6 ispitanika, a zadatak s participom *-(y)An* (B1–5) njih 5. Riječ je o zadacima u kojima su ispitanici morali razlikovati više različitih glagolskih oblika. Npr. u zadatku A2–5 za rečenicu „Rekao je da ne idem.“ bili su ponuđeni sljedeći odgovori: a) *Gitmedigimi söyledi*. b) *Gitmememi söyledi*. c) *Gitmeyeceğimi söyledi*. d) *Gitmemi söyledi*. Nekoliko različitih infinitnih oblika u ponuđenim odgovorima učinilo je taj zadatak izrazito zahtjevnim pa su bila samo 4 točna odgovora.

Međutim, kada u ponuđenim odgovorima ne variraju različite infinitne kategorije, nego se izmjenjuju isključivo posvojni sufiksi na infinitnoj formi, uspješnost se znatno povećava. To potvrđuju pitanja B1–6 i B1–7, gdje je broj točnih odgovora bio 11 odnosno 9. Primjerice, u zadatku B1–6 za rečenicu „Kad sam otvorila vrata, on je već odlazio.“ ispitanici su birali između ovih prijevoda: a) *Kapıyı açtığımda o zaten gidiyordu*. b) *Kapıyı açtığımda o zaten gidiyordu*. c) *Kapıyı açtığımızda o zaten gidiyordu*. d) *Kapıyı açtığımızda o zaten gidiyordu*. S obzirom na to da se točnost odgovora u takvim pitanjima temelji ponajprije na prepoznavanju posvojnog sufiksa, moglo bi se reći da bi takvi zadaci zapravo pripadali nižoj razini, ponajprije A1, gdje se posvojnost i obrađuje. Međutim, ta su pitanja namjerno uključena u razinu B1 kako bi se provjerilo na koji način na točnost utječe razlika u posvojnim sufiksima, odnosno hoće li ispitanici bolje rješavati one zadatke u kojima se ne mijenja infinitna forma, nego samo posvojni sufiks. Upravo zbog ta dva relativno jednostavnija zadatka, kao i zbog zadatka s kondicionalom, razina B1 pokazala je viši ukupni udio točnih odgovora (52,38 %) nego što bi se očekivalo s obzirom na složenost same razine. Međutim, na razini B2, u kojoj se posljednja tri pitanja odnose na složene infinitne forme, uspješnost je očekivano niska.

S obzirom na rečeno, istraživanje je potvrdilo početnu pretpostavku da će znanje infinitnih formi među ispitanicima biti znatno slabije od znanja finitnih glagolskih oblika. Na to nesumnjivo utječe činjenica da su infinitne forme područje koje govornici indoeuropskih jezika teže usvajaju, a i to što je njihova čestotnost u serijama, kako se u analizi korpusa moglo vidjeti, niža od čestotnosti finitnih formi. Kao ilustraciju za kraj možemo navesti da su dva najuspješnija ispitanika na pitanja odgovorila s tek po dva netočna odgovora, odnosno da su pitanja riješili s visokih 92,86 % točnosti. Njihovi netočni odgovori odnosili su se na pitanja A2–5 (glagolska imenica na *-mA*) i B2–7 (*-mAsI durumunda*), odnosno A2–6 (perfekt na *-mIş*) i B1–5 (particip na *-(y)An*). Dakle, i u slučajevima gdje poznavanje turskog jezika ide do visokih razina pogreške su se pojavile kod infinitnih oblika i perfekta na *-mIş*, što se uklapa u opći obrazac rezultata upitnika.

7. Zaključak

Cilj ovoga istraživanja bio je dokumentirati i barem djelomice opisati trend usvajanja turskoga jezika putem televizijskih serija – fenomen koji prije 2010. godine, kad je u Hrvatskoj prvi put započelo emitiranje jedne turske serije, nije bio ni zamisliv. Od tada do danas praćenje turskih serija, dodatno potaknuto internetskom dostupnošću sadržaja, postalo je stabilan i izrazito živ oblik kontakta hrvatskih gledatelja s turskim jezikom.

Ovaj je rad nastao kombinacijom pristupa primijenjene i korpusne lingvistike, s ciljem da se prvi put sustavno zabilježe jezični učinci izloženosti govornika hrvatskoga jezika turskim televizijskim serijama. Budući da je turski jezik kao neindeuropski sustav morfološki složen i za govornike hrvatskoga posebno zahtjevan, bilo je zanimljivo istražiti koje glagolske strukture gledatelji mogu spontano usvojiti isključivo putem serija. Iako se istraživanje moglo usmjeriti primjerice i prema pragmatičkim elementima poput pozdrava, čestih frazema ili formula ljubaznosti, ili npr. tako da obuhvati sve gramatičke sadržaje pojedinih jezičnih razina, fokus je stavljen na glagolske oblike jer oni čine temelj sintakse te omogućuju precizno vrednovanje i praćenje prema jezičnim razinama.

Korpusna analiza serije *Yabani* pokazala je da najfrekventniji glagolski oblici slijede potrebe televizijskog dijaloga. Najzastupljeniji su prezent glagola *imek*, predikativi *var* i *yok*, imperativ te perfekt na *-DI*, dok su složena vremena i infinitne kategorije znatno rjeđe. Takva raspodjela odražava strukturu žanra: brze replike, neposredna komunikacija među likovima i linearna radnja oslanjaju se na finitne glagolske oblike, odnosno na jednostavne rečenice.

Povezivanjem korpusne analize i analize podataka dobivenih upitnikom o namjernom usvajanju turskoga jezika moglo se vidjeti koji glagolski oblici prelaze u aktivni repertoar gledatelja koji jezik usvajaju isključivo putem medijskog sadržaja. Rezultati upitnika pokazuju da gledatelji spontano prepoznaju one oblike koji su i najfrekventniji u korpusu: finitne kategorije, prezentske i perfektne paradigme te jednostavne negacijske i upitne oblike. Znatno slabije prepoznati ostaju infinitni oblici i rijetke morfološke strukture, što je u skladu s njihovom manjom zastuplenošću u analiziranom materijalu.

Koliko nam je poznato, ovo je prvi rad usmjeren na ovaj aspekt gledanja turskih serija u Hrvatskoj. Bilo bi zanimljivo dodatnim istraživanjima potpunije opisati što gledatelji sve mogu naučiti, osobito u usporedbi s već brojnim radovima o usvajanju engleskoga jezika putem popularnih medija. Jedno od ključnih pitanja ovoga rada bilo je i to koliko je jezik serije u skladu sa standardnim jezikom, odnosno koliko je ta serija prikladan izvor za usvajanje suvremenoga turskog. Usporedba s rječnikom *A Frequency Dictionary of Turkish* pokazala je visok stupanj podudarnosti među najčešćim lemmama, što potvrđuje da je jezik serije uvelike usklađen s visokofrekventnim slojem standardnoga jezika i da se u tom pogledu može smatrati relevantnim izvorom.

Literatura

- Aksan, Y., Aksan, M., Koltuksuz, A., Sezer, T., Mersinli, U., Demirhan, U. U., Yilmazer, H., Atasoy, G., Öz, S., Yıldız, İ., i Kurtoglu, Ö. (2012). Construction of the Turkish National Corpus (TNC). U *Proceedings of the Eighth International Conference on Language Resources and Evaluation (LREC'12)* (str. 3223–3227). European Language Resources Association (ELRA).
- Aksan, Y., Aksan, M., Mersinli, Ü. i Demirhan, U. U. (2017). *A frequency dictionary of Turkish*. Routledge.
- Arslan, Ş. (2022). *Yumuşak güç algısında Türk televizyon dizileri: Turkish TV series in the perception of soft power* [doktorska disertacija, Sakarya Üniversitesi].
<https://acikerisim.sakarya.edu.tr/handle/20.500.12619/98550> (4.10.2025.)
- Bednarek, M. (2018). *Language and television series: A linguistic approach to TV dialogue*. Cambridge University Press.
- Bowker, L. i Pearson, J. (2002). *Working with specialized language: A practical guide to using corpora*. Routledge.
- Čaušević, E. (1996). *Gramatika suvremenoga turskog jezika*. Hrvatska sveučilišna naklada.
- Čaušević, E., i Kerovec, B. (2021). *Turski i hrvatski u usporedbi i kontrastiranju: Sintagma i jednostavna rečenica*. Ibis grafika.
- Çöltekin, Ç., Doğruöz, A. S. i Çetinoğlu, Ö. (2023). Resources for Turkish natural language processing: A critical survey. *Language Resources and Evaluation*, 57(2), 449–488.
<https://doi.org/10.1007/s10579-022-09605-4>
- d'Ydewalle, G., i Pavakanun, U. (1995). Acquisition of a second/foreign language by viewing a television program. U P. Winterhoff-Spurk (ur.), *Psychology of media in Europe: The state of the art – Perspectives for the future* (str. 51–65). Westdeutscher Verlag.
- Erdil, M. i Açıık, F. (2021). Investigation of grammar contents in Turkey Maarif Foundation's teaching Turkish as a foreign language program and in Turkish teaching sets. *International Journal of Languages' Education and Teaching*, 9(4), 129–161.
<https://doi.org/10.29228/ijlet.54059>
- Eryılmaz, R. (2024). Yabancı ve İkinci Dil Olarak Türkçe Öğretiminde Türk Dizi ve Filmleri: Bir Meta-sentez Çalışması. *TRT Akademi*, 9(22), 888–909. <https://doi.org/10.37679/trta.1519600>
- Gazi Üniversitesi TÖMER. (2016a). *Yabancılar için Türkçe dil bilgisi (A1–A2–B1)* (4. izd.). Başak Matbaacılık.
- Gazi Üniversitesi TÖMER. (2016b). *Yabancılar için Türkçe dil bilgisi (B2–C1)* (3. izd.). Başak Matbaacılık.
- Kilgarrieff, A., Baisa, V., Bušta, J., Jakubiček, M., Kovář, V., Michelfeit, J., Rychlý, P., & Suchomel, V. (2014). The Sketch Engine: Ten years on. *Lexicography*, 1(1), 7–36. <https://doi.org/10.1007/s40607-014-0009-9>
- Montero Perez, M., Sydorenko, T., i Huntley, L. (2024). Incidental vocabulary learning from watching audiovisual input. U M. F. Teng i B. L. Reynolds (ur.), *Researching incidental vocabulary learning in a second language* (str. 35–50).

- Nuroğlu, E. (2013). Dizi turizmi: Orta Doğu ve Balkanlar'dan gelen turistlerin Türkiye'yi ziyaret kararında Türk dizileri ne kadar etkili? U 5. *Uluslararası İstanbul İktisatçılar Zirvesi* (konferencijsko izlaganje).
- Ofłazer, K. i Saraçlar, M. (2018). Turkish and its challenges for language and speech processing. U K. Ofłazer i M. Saraçlar (ur.), *Turkish natural language processing* (str. 1–21). Springer. https://doi.org/10.1007/978-3-319-90165-7_1
- Podvezanec, T. (2025). *Računalna analiza i minimalni vokabular turskih dramskih serija na primjeru serije Yabani* [diplomski rad, Filozofski fakultet Sveučilišta u Zagrebu]. <https://repozitorij.ffzg.unizg.hr/en/object/ffzg:12504>
- Sadiç, F. (2024). Azerbaijani TV viewers' practices of watching Turkish television programs in the context of cultural studies: Turkish culture and tourism. *Erciyes İletişim Dergisi*, 11(2), 481–495. <https://doi.org/10.17680/erciyesiletisim.1444556>
- Sezgin, H. i Öztürk, M. S. (2020). A corpus analysis on the language on TV series. *Journal of Language and Linguistic Studies*, 16(1), 238–252. <https://doi.org/10.17263/jlls.712787>
- Teng, M. F. (2024). Introduction to researching incidental vocabulary learning in a second language. U M. F. Teng i B. L. Reynolds (ur.), *Researching incidental vocabulary learning in a second language* (str. 1–16).

Korpus

Yabani. Turska televizijska serija. Režija: Çağatay Tosun. Scenarij: Hilal Yıldız i Yekta Torun. NOW / NTC Medya, 2023–.

Digitalni alati

Sketch Engine. <https://www.sketchengine.eu/> (2003).

Verbal Forms in Turkish Television Series: Corpus Analysis as a Resource for Research on Incidental Language Acquisition

The aim of this paper was to document, and at least partially describe, the trend of acquiring Turkish through television series in Croatia – a phenomenon present since 2010, when a Turkish series was first broadcast on a domestic television channel. Since then, the consumption of Turkish series, further encouraged by the online availability of content, has become a stable and remarkably dynamic form of contact between Croatian viewers and the Turkish language.

This study examines the distribution of verbal forms in the first five episodes of the Turkish television series *Yabani* in order to assess the extent to which viewers may spontaneously acquire elements of the Turkish language through watching series. The corpus, consisting of 32,253 words, was processed using *Sketch Engine*, with necessary manual corrections due to the limited accuracy of existing morphological annotation tools for Turkish. A comparison with *A Frequency Dictionary of Turkish* showed that approximately 50 % of the most frequent lemmas overlap with the language of the series, while the overlap for the most frequent verbs reaches 74 %.

The corpus analysis extracted all verbal forms, calculated their frequencies, and established their relative proportions and distributions across categories. The most frequent forms in the corpus are the present of *imek* 'to be', the imperative, and the *-DI* perfect, while infinitival forms and complex tenses are considerably rarer. The analysis further indicates that singular forms predominate: first- and second-person forms due to the dialogic nature of the text, and third-person forms owing to the narrative passages.

The questionnaire conducted with viewers of Turkish series (N = 15) shows that participants most successfully recognise finite forms such as the present tense and the imperative, whereas recognition of non-finite and more complex forms is markedly lower. These results confirm the initial assumption that basic finite verbal structures are the most readily acquired through exposure to series.

Keywords: Turkish language; television series; corpus analysis; verbal forms; incidental language acquisition