

LA INTEL·LIGÈNCIA ARTIFICIAL A DEBAT

Una experiència de Iubilo URV

Isabel Baixeras, Rosa Queral, Carme Crespo,
Jaume Aymí, Jordi Blay i Albert Bordons (coord.)

LA INTEL·LIGÈNCIA ARTIFICIAL A DEBAT

Una experiència de Iubilo URV

Isabel Baixeras, Carme Crespo, Jordi Blay, Jaume Aymí,
Rosa Queral i Albert Bordons (coordinador)



Tarragona, 2025

PUBLICACIONS DE LA UNIVERSITAT ROVIRA I VIRGILI

Av. Catalunya, 35 - 43002 Tarragona

Tel. 977 558 474 · publicacions@urv.cat

www.publicacions.urv.cat



1a edició: setembre de 2025

ISBN (paper): 978-84-1365-236-8

ISBN (PDF): 978-84-1365-235-1

DOI: 10.17345/9788413652368

Dipòsit legal: T 907-2025



Cita el llibre.



Consulta el llibre a la nostra web.



Llibre sota una llicència Creative Commons BY-NC-SA.

Publicacions de la Universitat Rovira i Virgili és membre de la Unión de Editoriales Universitarias Españolas i de la Xarxa Vives, fet que garanteix la difusió i comercialització de les seves publicacions a nivell nacional i internacional.

CONTINGUTS

Agraïments	9
Pròleg	11
Prefaci	15
1. Introducció	19
2. La nostra experiència.....	35
3. Amenaces i riscos de la IA	45
4. Oportunitats i beneficis de la IA	55
5. Refutacions principals als arguments a favor i en contra de la IA	69
6. Conclusió: la IA és una amenaça o no?.....	81
7. Conclusions de la nostra experiència	87
8. Els relats de cadascú.....	91
9. Jubilo URV, o el valor d'una experiència sobre la jubilació: Compartint emocions	121
10. La Xarxa Vives d'Universitats i la Lliga de debat sènior	135
11. Breu glossari de termes de la IA.....	141
12. Preguntes a l'aplicació d'IA conversacional ChatGTP V2 ..	145
13. Bibliografia	149
14. Imatges de la competició.....	159

«La imaginació ens porta a mons que mai
van existir. Però sense ella, no anem enlloc.»

Carl Sagan

«I em queda secretament,
i no ho veureu,
la por
de l'espai sense res
i del desert que em capta»

Felícia Fuster, fragment del poema IX
del llibre *I encara*, Editorial 3i4, 1987.

AGRAÏMENTS

En primer lloc, donem les gràcies sobretot al company José Eugenio García-Albea, que va ser membre actiu de l'equip en tot moment, tot aportant les seves bones idees i suggeriments en la fase preparatòria i participant en la competició amb entusiasme i dedicació. També agraïm als companys Maria Rull i Estanis Pastor, que van ser membres de l'equip en la fase preparatòria, havent-hi col·laborat amb molt bones aportacions.

Donem el nostre agraïment als companys de Iubilo URV, que ens han donat suport constant en aquest repte.

Agraïm especialment a la Universitat Rovira i Virgili haver-nos donat l'oportunitat de participar en aquesta experiència magnífica i engrescadora, i en particular al vicerector de Compromís Social i Sostenibilitat, Jordi Diloli, que ens va encarregar a la comunitat Iubilo URV la participació en la Lliga de Debat Sènior del 2024.

Donem les gràcies a la Xarxa Vives d'Universitats per l'organització de la Lliga de Debat Sènior 2024, en la qual hem participat.

Agraïm a la Universitat Politècnica de Catalunya la seva hospitalitat en acollir-nos els dies de la competició, especialment en un moment difícil amb quasi plena ocupació dels seus locals per altres compromisos.

Ens sentim agraïts amb els companys de les universitats rivals a la Lliga de Debat Sènior, amb els quals vam competir, pel seu bon tarannà i afabilitat.

Agraïm al Servei de Publicacions de la Universitat Rovira i Virgili i al seu coordinador, Jaume Llambrich, la tasca d'aquesta publicació i els consells en la seva elaboració.

També ens complau reconèixer la bona feina i ajut de Iolanda Santiago, cap de l'Oficina d'Igualtat i Compromís Social.

Donem les gràcies en especial a la Diputació de Tarragona, pel suport en aquesta publicació.

També agraïm en particular a l'expert en projectes de comunicació Oriol Margalef per la bona revisió d'estil i format del text.

Igualment agraïm a Roger Mor Crespo i a la Isabel Oussedick Vilalta les imatges cedides.

I finalment, agraïm a les nostres respectives famílies la paciència i la comprensió al llarg de la preparació del debat, dels dies de la competició i durant l'elaboració d'aquest llibre.

PRÒLEG

La intel·ligència artificial generativa ha esdevingut un dels fenòmens més transformadors del nostre temps, una revolució tecnològica que sacseja les bases de la societat, l'economia i la cultura. Ens trobem davant d'una eina poderosa, amb un potencial il·limitat per millorar la nostra vida, però també amb la capacitat de plantejar reptes ètics, filosòfics i socials. Aquesta dualitat —entre oportunitat i amenaça, entre progrés i risc— ha estat el punt de partida del treball que teniu a les mans, fruit de la participació de l'equip de Jubilo URV a la Lliga de Debat Sènior de la Xarxa Vives d'Universitats.

Aquest llibre no és tan sols una recopilació d'arguments a favor i en contra de la pregunta central del debat —«La intel·ligència artificial és una amenaça per a la humanitat?»—, sinó també un testimoni del rigor, l'entusiasme i la capacitat de reflexió d'un grup de persones que, malgrat que han deixat enrere l'etapa professional activa, segueixen compromeses amb l'estudi, la discussió i la creació de coneixement.

Com a rector de la Universitat Rovira i Virgili sempre que en tinc ocasió m'agrada destacar el principal capital que tenim, que no és altre que les persones. L'altre gran actiu d'una universitat és el coneixement, que té un valor acumulatiu, de tal manera que com més gran és una persona, més valor adquirit posseeix. Per això és fonamental que aquells que han estat part de la institució durant molts anys no se'n desvinculin en el moment de la jubilació. L'experiència i el bagatge de generacions de docents, investigadors i professionals segueixen sent una font de riquesa per a la Universitat i per a la societat. I a la URV ens podem felicitar de tenir una associació com Jubilo URV, que amb la participació per primera vegada en la Lliga de Debat Sènior exemplifiquen aquest valor.

Més enllà d'explicar-nos en primera persona l'experiència, a través d'aquestes pàgines es posa de manifest la importància de mantenir actiu el pensament acadèmic en totes les etapes de la vida i el seu treball rigorós, que pot servir de model a generacions futures, tant pel que fa a la metodologia com també com a exemple de la importància dels arguments fonamentats per defensar posicionaments. I és que un dels valors de la Universitat és la formació de persones amb esperit crític, amb capacitat d'anàlisi i amb habilitats per argumentar de manera rigorosa, qualitats essencials no només en l'àmbit acadèmic, sinó també en la construcció d'una societat democràtica i informada. Aquest llibre és una mostra palpable d'aquesta funció i del compromís dels seus membres amb el diàleg i la reflexió col·lectiva.

Per això, més enllà de la pregunta sobre si la IA és una amenaça o no, el veritable aprenentatge d'aquest llibre rau en la necessitat d'abordar aquests desafiaments amb esperit crític, curiositat i compromís ètic. La tecnologia no és bona ni dolenta en si mateixa: depèn de nosaltres, com a societat, com la desenvolupem i com la fem servir. Convido, doncs, el lector a endinsar-se en aquesta experiència, a deixar-se interpel·lar pels arguments i a sumar-se al debat. Perquè, en darrer terme, la clau del futur de la intel·ligència artificial no es troba en les màquines, sinó en les decisions que prenguem com a societat i en els valors que decidim defensar.

Finalment, en nom de la institució que represento, aprofito aquest aparador per donar les gràcies i felicitar Iubilo URV. Vull expressar el meu agraïment més sincer a totes les persones que en formen part; la seva implicació constant amb la Universitat, l'entusiasme que hi posen i la dedicació diària ens han de servir d'inspiració. Gràcies per continuar enriquint el nostre entorn acadèmic i social amb la vostra experiència, el vostre coneixement i la vostra passió pel saber. També vull felicitar-vos públicament pel gran paper que va fer en aquesta edició de la Lliga de Debat Sènior; una felicitació extensiva a tot l'equip i una menció especial a Isabel Baixeras Delclòs, que va ser reconeguda com a millor oradora de la competició.

Només em queda esperonar-vos perquè repetiu l'experiència i continueu aportant coneixement a la societat. M'agradaria pensar que és gràcies al paper d'associacions com Iubilo, que potencien la longevitat activa i saludable, que cada cop som més a prop de convertir-nos en una d'aquestes anomenades zones blaves del planeta, que són regions en les quals l'esperança de vida és molt superior a la resta del món. Més enllà d'hàbits saludables, els estudis que intenten esbrinar-ne les causes determinen que tenir una xarxa social forta i mantenir la sensació de propòsit a la vida en són algunes de les claus. Que per molts anys Iubilo URV sigui una eina útil per exercir aquestes funcions.

JOSEP PALLARÈS MARZAL
Rector de la Universitat Rovira i Virgili

PREFACI

Quan Jubilo em proposà de fer una breu introducció en aquesta publicació, ultra sentir-me'n sincerament honorat i, doncs, agraït, vaig experimentar una emoció especial, per una raó principal: poder emparellar en un sol objectiu dues corporacions de què la realitat em dugué a ésser fundador: la URV, en primer lloc, en tant que primer rector, i l'Institut Joan Lluís Vives, de què em fou lliurat el diploma de rector fundador, que nasqué amb la voluntat de funcionar com una plataforma d'interacció global i permanent al servei del sistema universitari català, que en tenia necessitat peremptòria amb el creixement eminent entre 1991 i 1992, amb la creació de quatre noves universitats a Catalunya: UPF, UdG, UdL i URV.

El propòsit fonamental era establir una connexió d'ampli abast en tots els nivells dels àmbits principals: recerca, docència i administració i gestió; apostàrem fort per disposar d'un lligam transversal que permetés unes relacions fluides entre totes les que estem unides per una llengua, una cultura i una història comunes. Vull subratllar d'una manera particular l'entusiasme i la contribució decisiva en aquesta tasca dels rectors Nadal Batle (UIB), Carles Solà (UAB), Josep M. Nadal (UdG), Víctor Ciurana (UdL), Enric Argullol (UPF) i jo mateix (URV). Érem conscients de la necessitat d'un instrument que vertebrés un intercanvi estable i àgil per facilitar i incentivar l'eficàcia, la informació exhaustiva, la transparència, la col·laboració activa, el debat, la mobilitat... en el terreny de dimensions considerables que configuren les universitats que compartim la llengua catalana i, doncs, una determinada i peculiar idiosincràsia.

El 28 d'octubre de 1994 es constituí l'Institut Joan Lluís Vives a Morella amb 13 universitats adherides. En el punt de partença, ens hi trobarem les de Catalunya, la de les Illes Balears, la de Perpinyà Via Domitia i la Jaume I de Castelló; però no es trigà gaire a comp-

tar amb les de València, Politècnica de València, Alacant i, més endavant, la Miguel Hernández d'Elx i la Universitat d'Andorra. Avui formen la xarxa 22 universitats. Tot i que la concepció inicial es fixà prioritàriament en el caràcter públic de les institucions, d'una manera espontània i òbvia s'hi afegiren també algunes privades (UOC, URL, UIDC), entusiastes de la iniciativa: el quadre completíssim, per tant, del sistema d'ensenyament superior.

Allò que batejava en els nostres desigs era convertir la singularitat de cada institució, tot respectant-ne escrupolosament l'autonomia inviolable que li correspon, en una sola universitat dels Països Catalans; en un espai extens que reforcés, a través de la diversitat que la suma de cadascuna podia aportar, els llaços d'unió dels territoris que els formen.

No ens ha mogut mai, ni en els moments embrionaris ni després, cap voluntat provinciana ni molt menys de tancament, d'aïllament respecte a ningú. Vam voler reflectir el desig de projecció des del primer moment amb l'elecció del nom que havia de dur la Xarxa: Joan Lluís Vives, una personalitat insigne i universal que, des de la filosofia, des del pensament, sortí de València (1493), on va néixer i estudiar a la universitat, per recórrer Europa, fins al final de la seva vida (Bruges 1540). En un altre sentit, amb l'homenatge a la figura d'un dels més destacats humanistes del segle XVI, vam voler alhora posar en relleu la importància de les ciències humanes i socials en els processos de transformació positiva de la realitat mitjançant la intel·ligència i el coneixement; més encara en un món tan necessitat avui d'anàlisi i de canvis amb capteniments crítics, davant la situació gens engrescadora del món a les acaballes del primer quart del segle XXI, tan condicionat per un materialisme poc favorable a la solidaritat i a la cerca del benestar de tothom.

La Xarxa Vives, nova denominació amb què coneixem l'Institut Joan Lluís Vives, té com a funció, i penso que ha de continuar tenint-la, l'optimització dels recursos materials i humans en els diversos camps de la vida a les universitats; l'aprofitament de l'enorme enriquiment que suposa que cada universitari dels qui la formen no

ho és d'una de sola, sinó de 22. Ampliar, per tant, molt considerablement el territori d'influència directa.

Vull tancar aquestes reflexions potser massa òbvies posant l'accent en una iniciativa creada a la URV a través justament d'una altra plataforma destinada a intensificar i prosseguir, més enllà de les disposicions administratives a què ens hem de sotmetre, en allò que ha d'ésser essencial en qualsevol universitat: el debat, l'examen detingut des de posicions diferents sobre temes igualment variats; mantenir els lligams amb la URV; programar activitats de formació continuada. Em refereixo a Iubilo URV; el nom 'alegrar-se' ja anticipa una saludable resistència a com s'entén burocràticament la jubilació. No volem abocar-nos al final de viure intensament, al tancament de cap porta a fer créixer sense cap aturador i indefinidament les nostres capacitats; a compartir l'experiència llarga que atorga al pensament la consistència del treball quotidià, permanent, durant tants d'anys; a no desaprofitar absurdament els nivells assolits per l'experiència en un aïllament que lesionaria també les relacions socials entre col·legues.

Iubilo URV ha intervingut darrerament en una competició que ha organitzat la Xarxa Vives, com a joc de discussió entorn de la capacitat dialèctica entre sèniors respecte a un concepte que ara amonina especialment, el de la intel·ligència artificial. I dic expressament com a joc, perquè és un mitjà d'aprenentatge especial i principal en les fases inicials de la vida; demostrem, així, que jugar és una activitat útil de formació també en les etapes en què ens trobem els qui participem activament amb Iubilo. Em consta que Iubilo hi ha fet un paper brillant. Pot titllar-se d'infantilisme, però més igual, que expressi la il·lusió que sento amb els qui, des de posicions contrastades, van competir-hi. Enhorabona! La Xarxa ha fet d'escenari ideal per a la intervenció d'universitaris que ja no practiquen allò que han fet en llur vida laboral, però que fan una tasca excel·lent i profundament lloable, que desperta una alegria viva; viva; sempre viva!

JOAN MARTÍ I CASTELL
Rector fundador de la Xarxa Vives

1. INTRODUCCIÓ

Per començar, cinc cèntims. Aquí s'introdueix com uns quants «sèniors» de la URV han participat en uns debats sobre si la intel·ligència artificial (IA) pot ser una amenaça per a la humanitat o no i, sobretot, es fa un discerniment dels tres conceptes principals emprats: què s'entén per IA, per amenaça i per humanitat.

Els prejudicis i la por davant la intel·ligència artificial

«Ara m'he convertit en la mort, el destructor de mons».

Aquesta citació del **Bhagavad Gita**, un text fonamental de la filosofia hindú, va passar pel cap de J. Robert Oppenheimer, el 16 de juliol del 1945, mentre presenciava la primera detonació de la bomba atòmica, l'arma nuclear que ell havia contribuït a crear. Sembla que Oppenheimer va mantenir durant la segona meitat de la seva vida un doble pensament al voltant d'aquest enginy: orgull per la fita tècnica que suposava i pel desenvolupament d'una energia que havia de generar un gran benefici a la humanitat, però també culpa pels seus efectes sobre la població i el territori. Al final de la vida, ell pensava que l'energia atòmica era un problema perquè superava el context intel·lectual (i probablement polític i social) del seu temps (**Sadurní**, 2024).

Quan l'equip de la Universitat Rovira i Virgili (d'ara endavant, URV) vam començar a treballar el tema escollit per a la Lliga de Debat Sènior 2024 de la Xarxa Lluís Vives d'Universitats («La intel·ligència artificial és una amenaça per a la humanitat?») era inevitable remetre's al creador de la bomba atòmica, aleshores de moda a causa de l'oscaritzada pel·lícula *Oppenheimer*, de Christopher Nolan

(2023) que reviu la seva trajectòria. I és que ja des de l'inici la intel·ligència artificial (la IA, a partir d'ara) apareixia tant en els articles científics com en els de divulgació que anàvem llegint amb aquest doble vessant, positiu i negatiu. Positiu pel que pot representar —i de fet ja representa— d'avenç tècnic i de potencial benefici per a la humanitat, però també negatiu per les possibilitats perverses que ofereix en mans i activitats inadequades.

En aquest sentit, d'entrada hom pot pensar que comparar la IA amb l'energia atòmica (i especialment amb la bomba) és exagerat, però la lectura de segons quines publicacions i articles i l'anàlisi de segons quins documents audiovisuals fins i tot deixa la bomba atòmica més que curta davant la perspectiva d'un mal ús de la IA.

Visió esbiaixada

Els debats mediàtics exageren tant els perills com els beneficis de les noves tecnologies, i fomenten els prejudicis i la desinformació.

És clar que pot ser que els articles de divulgació que en fan veure els potencials efectes perversos siguin simplement el resultat del típic recurs comunicatiu de diaris i revistes que busquen poder posar algun titular ben atractiu (i catastròfic, si pot ser) per enganxar el màxim de lectors, mentre que el contingut és força més vague. I a l'inrevés, també pot ser que els escrits i documents audiovisuals que divulguen els beneficis reals o potencials de la IA n'exagerin els aspectes positius o en parlin com si fos una tecnologia actual quan en realitat s'exposen potencialitats només desenvolupables al cap d'uns quants decennis.

Però ja estem entrant en matèria... No ens avancem! Abans de fer veure els arguments a favor o en contra de la IA, hem d'explicar de què va tot això de la Lliga de Debat Sènior de la Xarxa Lluís Vi-ves, i com va ser que un grup de jubilats actius de la URV va decidir participar-hi.

La URV participa en el debat sènior de la Xarxa Lluís Vives

La primavera del 2024, la Universitat Rovira i Virgili ha participat per primera vegada en la Lliga de Debat Sènior de la Xarxa Vives d'Universitats. Al capítol 10 d'aquest llibre comentem què és la Xarxa Vives d'Universitats, quines funcions té, quins debats organitza anualment i què és, en concret, la Lliga de Debat Sènior. La URV s'hi ha presentat amb un equip format íntegrament per persones que pertanyen a la comunitat Iubilo URV: un conjunt de membres de la URV que, tot i que ja s'han retirat, segueixen tenint iniciativa, coneixement i talent, individual i col·lectiu, i tot plegat ho posen al servei de la societat. Vegeu, al capítol 9, l'origen, les característiques, els objectius i activitats de Iubilo URV amb la visió personal de la coordinadora, Rosa Queral.

L'equip

Tot just concebut el projecte de la URV d'incorporar-se a la Lliga de Debat, una pluja d'idees practicada en el si de la Junta Executiva de Iubilo URV que coordinen Rosa Queral i Jaume Aymí va suggerir els noms de les persones associades que tal vegada podrien tenir la disposició de formar part de l'equip. El Jaume, que va ser l'encarregat de convidar-los-hi, va haver de desplegar tots els seus dots de convenciment per vèncer la resistència inicial: massa complex, massa compromís, massa feina... Feliçment, però, va aconseguir formar l'equip que ell mateix havia de capitanejar durant tot el procés.

A més de Jaume Aymí (professor jubilat de Didàctica de la Llengua, expert en filologia catalana i en cinema, Facultat d'Educació i Psicologia, URV), ens hi vam anar incorporant Isabel Baixeras (advocada jubilada, experta en dret administratiu i síndica de greuges de la URV en el període 1999-2004), Estanislau Pastor (catedràtic emèrit de Psicologia Evolutiva i de l'Educació, expert en psicologia infantil i orientació psicopedagògica, Facultat d'Educació i Psicologia, URV), Albert Bordons (catedràtic emèrit de Bioquímica i

Biotecnologia, expert en microbiologia d'aliments, Facultat d'Enologia, URV), Carme Crespo (química jubilada, experta en grans equips analítics del Servei de Recursos Científics URV i exregidora de l'Ajuntament de Tarragona), José Eugenio García-Albea (catedràtic emèrit de Psicologia, expert en psicolingüística i psicologia cognitiva, Facultat de Ciències de l'Educació i Psicologia, URV) i Maria Rull (professora jubilada d'Anestesiologia, experta en tractament del dolor crònic, Facultat de Medicina i Ciències de la Salut, URV). Unes setmanes més tard, quan l'Estanis i la Maria, absorbits per altres compromisos, ho van haver de deixar, s'hi va afegir Jordi Blay (professor jubilat de Geografia, expert en medi rural, Facultat de Turisme i Geografia, URV). L'entesa entre nosaltres va ser immediata, i de seguida vam aconseguir la cohesió que l'equip necessitava i les ganes de passar-ho bé, tant en la preparació com en la celebració del debat.

Motivació clau

Les ganes de passar-ho bé van imposar-se a la diversitat d'orígens i a la falta d'experiència de treball en comú. L'equip de debat de Jubilo URV va assolir de seguida la cohesió necessària.

L'objecte del debat

La Lliga de Debat Sènior de la Xarxa Vives d'Universitats sempre es produeix al voltant d'una pregunta que cal respondre, de manera argumentada, en un sentit o altre; és a dir, a favor, o bé en contra, en diferents confrontacions per parelles entre els grups de les diverses universitats. Un sorteig que es fa uns moments abans de començar cada sessió decideix el posicionament de cada equip. La pregunta de la convocatòria del 2024 era «La intel·ligència artificial és una amenaça per a la humanitat?». Les posicions que podrien ser-nos adjudicades en els successius debats de la Lliga serien l'afirmativa (que sí, que la intel·ligència artificial és una amenaça per a la huma-

nitat) o bé la negativa (que no, que la intel·ligència artificial no és cap amenaça per a la humanitat).

El tema va despertar un gran interès entre els membres de l'equip. És una qüestió d'actualitat i de futur. Ja fa uns quants anys que la nova tecnologia coneguda com a intel·ligència artificial ha entrat en les nostres vides, i cada vegada hi és més present. Creix d'una manera vertiginosa, i la seva expansió exponencial provoca la desconfiança que causa tot allò que no es coneix. L'alerta davant un eventual descontrol, la por davant l'acumulació de poder i la constatació dels riscos en relació amb la pervivència de les construccions i organitzacions humanes aconseguides fins ara són alarmes capaces de qüestionar la conveniència i l'oportunitat dels beneficis que aquest recurs ja ens reporta, avui dia, en el progrés de la ciència, la medicina, la tècnica, l'administració dels afers comuns, els entreteniments...

La preparació per al debat ens va requerir un doble esforç: per una banda, una major obertura de ment per no rebutjar allò que és nou només perquè sigui desconegut; per altra banda, una reflexió sobre els valors assolits fins ara per les diferents civilitzacions i cultures de la Terra; uns valors que no s'haurien d'extraviar en l'ambigua i confusa barreja d'idees, conceptes i llenguatges que la IA ja està confegint.

Per preparar les nostres intervencions en el debat havíem d'aclarir, d'entrada, els conceptes principals que hi hauríem de manejar: *intel·ligència artificial*, *amença* i *humanitat*. Seguidament, hauríem d'abordar la selecció i la classificació de les fonts que ens podrien ser útils en la lliga: antecedents filosòfics, textos doctrinals, articles científics i tecnològics, opinions periodístiques, imatges fotogràfiques, mitjans audiovisuals, etc. En una tercera etapa, hauríem de construir, amb les idees recollides, una col·lecció d'esquemes o esborranys de discursos diferents —tant en un sentit com en l'altre, és a dir, tant a favor com en contra de la proposició— a fi de tenir a la nostra disposició una bona borsa d'arguments preparats per esgrimir-los a les ponències i a les rèpliques. Finalment, hauríem

d'assajar les exposicions del debat previstes en el reglament de la lliga (introducció, ponència, rèpliques i conclusions) i fer-nos conscients que haurien de ser ajustades al cronòmetre, lingüísticament correctes, assertives en la forma, convincents en el fons, considerades amb el contrincant i polides amb el jurat.

El resultat de la recerca

Els membres de Iubilo URV no han pas perdut, pel fet de ser grans, el talent que els va fer universitaris. Conserven l'hàbit d'estudiar, la capacitat de reflexionar i l'habilitat d'investigar que van adquirir durant la seva vida activa a la Universitat. És per això que es pot dir, sense pecar d'immodestos, que el resultat del treball de vuit persones universitàries durant nou o deu setmanes del 2024 no només ha proporcionat molta satisfacció als participants sinó que la seva sistematització ha generat coneixement. El resultat de la recerca pertany a la Universitat Rovira i Virgili —l'entitat de la qual Iubilo URV forma part—, i aquesta publicació del Servei de Publicacions URV el posa a disposició de qui hi pugui estar interessat.

El talent no té edat

«Els membres de Iubilo-URV no han perdut el talent que els va fer universitaris. Conserven l'hàbit d'estudiar, la capacitat de reflexionar i l'habilitat d'investigar.»

Cert és que la tecnologia de la IA creix a un ritme exponencial, i que en aquesta matèria allò que avui és d'actualitat demà podria estar superat. Som molt conscients que els riscos que ara a principis del 2025 s'atribueixen a la IA aviat seran obsolets. Però sabem també que la ciència avança amb passos vacil·lants, i que totes les reflexions, per petites que siguin, poden contribuir a la formació de la consciència global.

És per això que no només compartim, seguidament, el procés de formació dels criteris utilitzats a la Lliga de Debat, sinó també el resultat de la nostra recerca, és a dir, els principals arguments

que propicien sostenir que la intel·ligència artificial és realment una amenaça per a la humanitat, i els arguments que descarten que els riscos anunciats derivin en una veritable amenaça. Tant de bo aquesta cerca feta amb caràcter científic constitueixi una petita transferència de coneixements per part nostra a la població en general.

Totes les referències tenen un suport bibliogràfic accessible i consultable amb hipertext a la versió digital i amb la llista alfabètica de referències bibliogràfiques, al capítol 13, junt amb un QR per accedir directament a les fonts des de la versió impresa.

Discerniment dels conceptes emprats en la proposició

El tema a debatre donava per fet que la humanitat té assumits uns objectius concrets i que, en l'acompliment d'aquests objectius, es podria haver d'enfrontar amb algunes dificultats originades per la intel·ligència artificial o generades en ocasió de la utilització d'eines d'intel·ligència artificial, que podrien arribar a ser conceptualitzades com una amenaça per a la humanitat.

Encara que la concreció dels tres conceptes (IA, amenaça, humanitat) podria haver estat passada per alt en una conversa de cafè, el prestigi de la comunitat Iubilo URV, el caràcter universitari de la Lliga de Debat i, sobretot, les ganes d'aprofundir en uns temes poc coneguts ens convidaven a immersir-nos-hi: conèixer el vocabulari de la intel·ligència artificial, avesar-nos al llenguatge que li és propi, aprendre a distingir les diferents eines, actualitzar constantment la informació... A més, havíem de discernir què és, ben bé, una *amença*, i ens havíem de posar d'acord en relació amb què vol dir *humanitat*, per emprar tots aquests termes amb prou naturalitat i, en la mesura que fos possible, amb una certa propietat.

Heus ací un resum de la petita recerca prèvia.

Què és la intel·ligència artificial

La intel·ligència artificial (també coneguda amb la sigla IA, o AI quan es diu en anglès) és el camp que implica màquines —o sigui, ordinadors— capaces d'imitar determinades funcions de la intel·ligència humana, incloses característiques com la percepció, l'aprenentatge, el raonament, la resolució de problemes, la interacció lingüística i fins i tot la producció de treballs creatius. A mesura que avançàvem en l'estudi del tema, i a partir de l'examen de diferents fonts, vam anar elaborant un *glossari* de termes relacionats amb la intel·ligència artificial, del qual hem fet un resum que trobem al capítol 11.

Existeix un cert consens que la intel·ligència humana és la capacitat mental que implica l'habilitat de raonar, planejar, resoldre problemes, pensar de manera abstracta, comprendre idees complexes i aprendre ràpidament, i de l'experiència; tot plegat, referit a la capacitat de comprendre l'entorn (Gottfredson, 1997). Es coneix com a intel·ligència artificial la part de la informàtica que desenvolupa algorismes —o sigui, uns conjunts d'instruccions per executar una tasca o resoldre un problema— per dotar un computador de capacitats de raonament que imiten alguns aspectes de la intel·ligència humana.

Les funcions de la intel·ligència que més interessin la IA són la percepció visual o auditiva, l'aprenentatge, la resolució de problemes, el suport a la robòtica, el processament de llenguatges naturals i la interacció lingüística. Interessa, també, la seva potencialitat en la producció de treballs creatius. Per assolir aquests objectius, els investigadors de la IA han adaptat i integrat una àmplia gamma de tècniques, incloent-hi la cerca i l'optimització matemàtica, la lògica formal, els mètodes basats en estadístiques, investigacions d'operacions i economia, i sobretot el desenvolupament de les xarxes neuronals, que tot seguit comentarem.

La importància de la IA i el seu boom actual deuen haver influït en el fet que els guardonats amb els Premis Nobel 2024 de Física i de Química hagin estat justament científics pioners en les bases

i les aplicacions de la IA. En concret, el de Física ha estat per als professors emèrits John Hopfield —de la Universitat de Princeton (EUA)— i Geoffrey Hinton —de la de Toronto (Canadà)—, pels descobriments sobre els fonaments de les xarxes neuronals i l'aprenentatge automàtic (*machine learning*). I el Nobel de Química ha estat per a Demis Hassabis i John Jumper, ambdós de l'empresa Deep Mind, amb el professor David Baker, de la Universitat de Washington (EUA), per les seves prediccions en l'estructura i el disseny computacional de proteïnes (Nobel Prizes, 2024).

El desenvolupament de la IA no es basa només en les matemàtiques i la informàtica, sinó que sovint recorre a disciplines com la psicologia, la lingüística, la filosofia, la neurociència i altres camps.

La lògica i la matemàtica tenen l'origen a les civilitzacions antigues, i es desenvolupen d'una manera o altra en totes les cultures de la Terra. Pel que fa a la nostra, en els treballs de Ramon Llull al segle XIII, en els de Pascal i Leibniz al segle XVII i en els autòmats de Torres Quevedo a principis del segle XX trobem alguns antecedents dels models computacionals que l'any 1936 Alan Turing va presentar en el camp que ell va anomenar intel·ligència de màquines. El film *Desxifrant l'enigma*, dirigit per Morten Tyldum el 2014, és un biopic sobre aquest brillant matemàtic anglès, que va desxifrar els codis secrets emprats pels nazis.

A finals dels anys 1950, la intel·ligència artificial es va establir com a disciplina acadèmica en algunes universitats britàniques i dels Estats Units; John McCarthy, Marvin Minsky, Nathaniel Rochester i Claude Shannon se'n consideren els fundadors.

Els anys següents la IA va passar per múltiples cicles d'optimisme, seguits d'uns períodes de decepció (amb la consegüent pèrdua de finançament) que es coneixen com l'hivern de la IA. Als anys 1990 l'expansió dels ordinadors personals i l'accés general a internet van revolucionar la informàtica i són la base de l'estat de la IA en l'actualitat. Una fita important va ser el supercomputador Deep Blue d'IBM, especialitzat a jugar a escacs, que el 1997 va aconseguir guanyar al campió mundial Garri Kaspàrov. Era el primer cop que

una màquina amb una IA molt específica superava els humans en operacions complexes. En qualsevol cas, ja feia molt de temps que qualsevol calculadora electrònica podia fer operacions matemàtiques bàsiques més ràpidament que els humans.

A principis del segle XXI es comencen a desenvolupar molts programes informàtics basats en el raonament deductiu logicomatemàtic que s'apliquen a sectors diversos, com ara les finances, la diagnosi mèdica, la recerca i els jocs —com Deep Blue. Es tracta de la *IA simbòlica*, en la qual els algorismes que els programadors han elaborat tenen regles explícites de formulació.

Tanmateix, la IA actualment més coneguda és la *IA no simbòlica*, o de connexions, que s'ha anat desenvolupant els últims 20 anys, i que es basa en mètodes inductius estadístics. En aquest cas, la clau són les *xarxes neuronals* artificials: uns models matemàtics d'aprenentatge automàtic (*machine learning*) inspirats en el sistema nerviós biològic, amb nòduls interconnectats. L'aprenentatge automàtic és la ciència de programar els computadors de tal manera que ells mateixos aprenguin de les dades que manipulen (Niazian & Niedbala, 2020). Aquest aprenentatge es basa en uns algorismes que processen exemples d'un conjunt enorme de dades. Les respostes que hi són possibles van canviant gradualment en repeticions successives a mesura que els encerts augmenten, de tal manera que el sistema aprèn a classificar les dades d'entrada dins d'un conjunt fix de possibles categories de sortida (Caudai et al., 2021). Quan l'aprenentatge no és supervisat, es coneix com a aprenentatge profund o *deep learning*.

Cap al 2012 comença el desenvolupament de les xarxes neuronals d'aprenentatge profund en multicapa, connectades l'una amb l'altra. Aquest desenvolupament del programari ha anat paral·lel a una gran millora en l'arquitectura dels transformadors, amb un augment de l'eficàcia del maquinari, que són els *supercomputadors*, sobretot després del 2017. Tot plegat ha donat lloc al boom del 2020: moltes empreses, universitats i laboratoris, sobretot a Califòrnia, s'hi dediquen intensament amb eines i aplicacions utilitzades globalment.

Les eines més conegudes d'aquest camp són les d'IA *generativa* i creativa de textos i imatges, però també ho són moltes aplicacions d'ús habitual, com els motors avançats de cerca a internet, els sistemes de recomanació i publicitat, la interacció amb la parla humana, la traducció automàtica de textos i veu, la utilització de mapes, els assistents virtuals, els bots de conversa, els cotxes autònoms i tants d'altres. De fet, moltes aplicacions d'IA d'avantguarda no es perceben com a IA i, sovint, quan alguna d'elles s'ha fet prou útil i habitual, ja no s'etiqueta com a IA. La IA actual —i no només la IA generativa— també s'aplica en molts sectors econòmics, incloses les administracions, les organitzacions d'educació, la comunicació, les indústries de tot tipus i, sobretot, la recerca i les seves aplicacions. Per a una major comprensió del funcionament de la IA, en general, recomanem el llibre de [Melanie Mitchell \(2024\)](#).

L'ús creixent de la intel·ligència artificial al segle XXI està influïnt en un canvi social i econòmic cap a l'augment de l'automatització, la presa de decisions basada en les dades i la integració dels sistemes d'IA en diversos sectors econòmics i àrees de vida. Aquest ús afecta els mercats laborals, la salut, el govern, la indústria, l'educació, la propaganda i la informació. Això planteja qüestions sobre els efectes de la IA a llarg termini, les seves implicacions ètiques i els riscos que genera, i provoca debats sobre les polítiques reguladores que seran necessàries per garantir la seguretat i els beneficis de la tecnologia.

Tots aquests possibles beneficis i riscos de la IA formen el contingut de les nostres confrontacions dialèctiques a la competició de la Lliga de Debat de la Xarxa Vives, tal com comentem en els altres capítols d'aquest llibre.

Què s'entén per amenaça

En la nostra preparació per al debat no podíem seguir endavant sense tenir clara quina és la diferència entre *amença* i *risc*. Una amenaça és un indicatiu pel qual hom manifesta un perill, una cosa

a témer, mentre que un risc és una contingència desfavorable a la qual està exposat algú o alguna cosa: un perill incert.

Com que la Xarxa Vives d'Universitats no ens va convidar pas a debatre si la IA comporta, o no, algun risc, sinó que el tema que ens plantejava era si vèiem aquesta innovació com una amenaça —és a dir, com un perill cert— calia establir, d'entrada, si les conseqüències de l'ús de la IA que ja es detecten en el moment actual es poden interpretar com a indicis reals, manifestacions clares d'un perill que ja existeix, o bé es tracta només d'aspectes potencials: un risc que encara no es manifesta en indicis preocupants.

Per poder considerar els impactes de la IA com a indicatiu preocupant (i, per tant, com a amenaça) cal tenir clar que el *risc* depèn de tres grans factors ([Institut de Seguretat Pública de Catalunya](#), 2017):

- Perillositat, composta de severitat o intensitat i probabilitat d'ocurrència d'un fenomen.
- Exposició, o quantificació de la ubicació d'un element o sistema amb relació a un perill concret, de manera que el fa vulnerable o susceptible de patir danys (nombre de persones i béns potencialment exposats).
- Vulnerabilitat, o predisposició intrínseca d'un element (sigui persona, edifici, organització, sistema, ecosistema, municipi, etc.) a patir danys per motiu d'un fenomen d'una intensitat concreta, és a dir, davant un perill concret.

Sovint s'expressa com a fórmula:

$$\text{Risc} = \text{Perillositat} \times \text{Vulnerabilitat} \times \text{Exposició}$$

Si no hi ha els tres elements, no hi ha risc.

Amenaça vs. Risc

Una amenaça implica un perill clar i existent, mentre que un risc representa una possibilitat desfavorable encara incerta.

Com cal prevenir el risc (o l'amenaça, si determinem que és així) que suposarà l'ús creixent de la IA? Les persones o entitats que tinguin possibilitats —o responsabilitats— d'actuació n'han de ser conscients i han d'impulsar estratègies de mitigació o reducció, tot actuant sobre un o diversos factors:

- Les estratègies antiperillositat, que implicarien prohibir o delimitar clarament l'ús de la IA en determinats camps que es considerin sensibles.
- Les estratègies antivulnerabilitat, que consistirien en la implantació d'eines de control efectiu, com per exemple uns programes capaços de detectar usos o impactes de la IA no permesos o indesitjables.
- Les estratègies antiexposició, que se centrarien en la informació i l'educació de la població enfront dels perills de l'ús de la IA.

Hi haurà *amença* per a la humanitat, doncs:

- Si el grau de perillositat que comporta l'ús de la IA pot afectar de forma notable la població i hi ha un mínim de probabilitat d'ocurrència, i no hi ha una regulació adequada.
- Si el nombre de persones i de béns potencialment afectats és tan ampli que no hi pot haver prou sistemes de control.
- Si la majoria de la població no està prou protegida davant els impactes de la IA perquè en desconeixen els efectes perversos en major o menor mesura.
- Si existeixen indicis —manifestacions clares— que això ja s'està començant a produir.

Per reforçar els nostres punts de vista, i tot seguint la indicació amable i experta de l'amic Juanjo Pujadas —l'antropòleg prestigiós que també és membre de Iubilo URV—, hem llegit o rellegit les tesis sobre la societat del risc del sociòleg Ulrich Beck ([Beck, 2006](#)). També pel consell del Dr. Pujadas hem accedit a la tesi doctoral de

la companya Rosa Queral Casanova, coordinadora de Iubilo URV, que tracta àmpliament les tesis de Beck en relació amb el risc.

Ens hem fet conscients del caràcter global dels riscos actuals, i de la seva condició moral relacionada amb un nou cosmopolitisme que coincideix a propugnar un major respecte envers el medi, els drets humans i la igualtat ètnica i de gènere. Amb U. Beck hem après que la socialització del risc —és a dir, el fet que diverses comunitats comparteixin un mateix risc— no és necessàriament un fenomen negatiu. D'una banda, el fet de ser-ne tots conscients genera responsabilitats noves. D'altra banda, per afrontar-lo cal superar les fronteres territorials.

La compartició global dels riscos de la IA podria ser, precisament, l'antítesi de l'amenaça: la necessitat d'afrontar-los per evitar-ne les conseqüències desfavorables podria ser la gènesi d'una nova comunitat política útil a la humanitat per a la consecució dels seus objectius.

Què és la humanitat i quins són els seus objectius

Per establir si la intel·ligència artificial destorbarà la humanitat en la consecució de la seva finalitat —i, per tant, per poder analitzar si la IA li pot arribar a ser una amenaça—, ens hem de preguntar què és concretament allò que la humanitat persegueix.

En la història de la filosofia el tractament del concepte *humanitat* és molt extens, però no tan divers com sembla. En qualsevol cas, la tesi d'Aristòtil referida a l'existència humana com un conjunt, la polis —que treballa per a la consecució d'un bé: el benestar de tot-hom, o la vida bona— no ha perdut vigència. Efectivament, 2.400 anys més tard, l'Agenda 2030 per al Desenvolupament Sostenible adoptada el 25 de setembre del 2015 per l'Assemblea General de l'Organització de Nacions Unides fixa, en uns termes anàlegs, els reptes immediats de la humanitat (almenys, els reptes immediats de la humanitat organitzada, representada d'alguna manera per l'organització internacional esmentada), per enfortir la pau universal i

l'accés a la justícia. Els ODS (objectius de desenvolupament sostenible de l'ONU (2015) tenen com a objecte les persones, el planeta i la prosperitat. L'estratègia concebuda és de combatre les desigualtats dins dels països, i també entre ells; de construir societats pacífiques, justes i inclusives, i de protegir els drets humans. Es tracta de trobar entre tots un major benestar col·lectiu. S'assembla molt a l'objectiu de «vida bona» que predicava Aristòtil, i no s'allunya gaire de la filosofia de la Pachamama i «el buen vivir» que impregna algunes de les noves constitucions dels països de l'Amèrica del Sud.

Aquest és el concepte d'humanitat que vam prendre com a referència per establir si, per a ella, la IA és una amenaça.

2. LA NOSTRA EXPERIÈNCIA

De la formació del grup a la construcció dels arguments, passant per la gestió del treball i el repartiment de feines. Aquí s'expliquen els detalls de la preparació que vam seguir els membres de Iubilo URV per participar en el debat.

La primera experiència directa de Iubilo URV a la Lliga de Debat Sènior que organitza la Xarxa Vives va ser observar per part del Jaume el funcionament real d'aquesta activitat el juny del 2023 a Elx, on va tenir lloc la fase final de la segona edició. En aquest cas es tractava del tema «L'humor té límits?». Tot molt pautat pel que fa a la posició que cada equip ha de defensar segons sorteig, a la limitació rigorosa del temps en cadascuna de les intervencions i especialment a l'esperit de respecte i fins i tot de cordialitat entre els diversos equips participants. Vet aquí un exercici d'intel·ligència col·lectiva, d'autèntica competició entre persones educades, amb les úniques armes de la paraula i l'argument i sense cap premi material que calgui aconseguir.

Després d'obtenir l'autorització de la Xarxa Vives d'Universitats per participar en la Lliga de Debat Sènior 2024, Iubilo URV es va disposar a formar l'equip que, en representació de la Universitat Rovira i Virgili, hi debatria sobre «La intel·ligència artificial és una amenaça per a la humanitat?». Els passos següents van ser els de formar l'equip, establir un calendari de reunions, buscar la informació, estudiar el tema i aconseguir l'entrenament necessari per a la fase de competició.

El mètode de treball va ser determinat pels mateixos membres de l'equip en les primeres sessions, i es va anar adaptant successivament a les necessitats que s'anaven fent evidents. Certament, el mètode hauria pogut ser qualsevol altre. De fet, algunes universitats preparen els equips amb la intervenció de personal especialitzat. Totes les tècniques poden ser bones, i totes tenen avantatges i inconvenients. Ara bé: l'autogestió de l'equip de Iubilo URV 2024 va ser un dels al·licients principals de l'experiència.

En qualsevol cas, amb la pretensió de ser útils als qui participin en edicions posteriors de la Lliga de Debat, sigui per adoptar, sigui per descartar la nostra tècnica, exposem seguidament les decisions adoptades per l'equip en matèria de mètode.

Detalls de mètode preparant el debat

CALENDARI DE REUNIONS, SISTEMATITZACIÓ I EMMAGATZEMATGE

Quan vam començar a treballar, al març del 2024, vam establir una dinàmica força clara: reunions setmanals del grup, tasques individuals a través de l'estudi i l'elaboració de fitxes i informes que es dipositaven en un espai comú de Google Drive creat per l'Albert, que mantenia les dades a l'abast de tothom, i comunicació interna en el dia a dia a través d'un grup de WhatsApp.

FORMACIÓ EN MATÈRIA DE COMUNICACIÓ

Abans de res, i per agafar impuls, vam participar en el taller virtual del Fòrum Vives titulat «Comunicació eficaç i creativa», que impartia l'experta en comunicació Fani Grande. Hi vam aprendre alguns dels aspectes més importants a tenir en compte en els debats dialèctics, com són la gestió emocional, la comunicació verbal i no verbal, la creativitat per organitzar els continguts dels parlaments, la posada en escena i altres. Va quedar clar que havíem d'allunyar les creences individuals (tot discernint creença i opinió, i distin-

gint opinió d'informació), ens havíem d'entrenar en la gimnàstica mental, havíem de practicar les habilitats comunicatives potencials i havíem de desplegar els recursos d'habilitat cognitiva que ens haurien de permetre defensar tant una argumentació a favor com una argumentació en contra del pronunciament que centrava el debat. Amb la mateixa intenció de preescalfament, el Jaume ens va proporcionar l'accés a l'estimulant documental *Asimov, un missatge per al futur*, de [Mathias Théry \(2022\)](#), que ens va omplir d'optimisme i de ganes de continuar l'estudi del tema.

Entrenar habilitats

«Va quedar clar que havíem d'allunyar les creences individuals, fer gimnàstica mental, practicar habilitats comunicatives i desplegar recursos d'habilitat cognitiva.»

NOCIONS D'INICIACIÓ

Al principi la majoria de nosaltres no tenia gaire idea de què era la IA, ni de les aplicacions ni de l'enrenou que suposava de cara al futur. José Eugenio era qui més hi entenia. En coneixia els fonaments i els orígens, ens parlava de les analogies entre la IA i la intel·ligència humana, i en destacava les diferències. Des del primer dia ens va parlar dels LLM (*large language models*) com a base de l'aplicació d'IA generativa conversacional ChatGPT, quan els altres encara no en sabíem gens.

En la reflexió sobre què s'entén per *humanitat*, l'Estanislau es referia a la diferència entre humans i altres animals: els humans han interposat la tècnica entre ells i el món, i els altres animals transformen molt menys el món.

En les primeres reunions ens vam aproximar al tema mitjançant les informacions o notícies que trobàvem als mitjans de premsa digital. El José Eugenio aportava sovint articles de premsa relacionats amb el tema. N'hi havia tant sobre els beneficis, usos, novetats i

progressos de la IA com sobre els seus possibles riscos i amenaces. Al llarg d'aquests mesos de treball, el bombardeig de notícies als mitjans de comunicació seria constant.

RÈGIM DE REUNIONS

L'equip es va preparar per a la competició en onze sessions presencials de treball en comú a la seu de Iubilo URV, al campus Catalunya de la URV. Se'n van aixecar unes actes que recollien els principals punts tractats i les tasques a fer. A banda, cadascun dels membres de l'equip desplegava quotidianament el treball individual de recerca i sistematització que havia assumit, per exposar-lo en la reunió següent.

DIVERSIFICACIÓ DEL TREBALL

La consolidació del grup va ser força ràpida. Tots llegíem i estudiàvem, però cadascú es va especialitzar en una part de la feina conjunta: específicament, l'Albert organitzava les informacions tecnològiques principals i elaborava i emmagatzemava els resums de cada sessió; la Maria aportava aspectes de recerca en el món de la medicina; l'Estanislau enriqueix el debat amb conceptes i valors filosòfics; la Isabel oferia passió i profunditat jurídica en l'argumentació; el Jaume presentava informacions audiovisuals de suport; la Carme estudiava el reglament de la Lliga, produïa fitxes de referència per a les discussions internes i espremia els seus contactes familiars amb experts en IA; el José Eugenio aportava solidesa en la teoria del coneixement i de l'aprenentatge i ens advertia si perdiem el fil; el Jordi suggeria idees noves i reconduïa alguns punts un pèl embussats. Si aconseguíem un bon clima de col·laboració i d'amistat, de tots aquests vímetes podia sortir un bon cistell...

COMPREENSIÓ DE LA PREGUNTA

Abans de respondre qualsevol pregunta és important haver-se assegurat que s'ha comprès bé quina és la qüestió plantejada.

Hem de reconèixer que, al començament, ens abordava sovint el dubte de si qualsevol argument o posició era *a favor* o *en contra* del repte plantejat, ja que instintivament —i erròniament— ho referíem a la mateixa intel·ligència artificial. Ens semblava que estar-hi a favor era reconèixer-li els beneficis que ens reporta i que estar-hi en contra era destacar els riscos i perjudicis que ens causa. Aviat ens vam centrar i vam ser conscients que la pregunta del debat era si «la IA és una amenaça per a la humanitat». Per tant, qui estigués a favor d'aquesta proposició, o sigui, qui donés una resposta afirmativa a la pregunta, respondria que sí, que és una amenaça per raó dels riscos que presenta. Aquesta reticència ve a ser un cert posicionament en contra de la IA. L'altra posició, o sigui, la de defensar que la IA no és cap amenaça per a la humanitat, implica que s'està en contra del plantejament, que se'n rebutja l'enunciat; o sigui, que s'està a favor de la IA, tot pensant en els seus beneficis i usos. Un bon embolic!

CONSCIÈNCIA DEL REGLAMENT

Vam tenir sempre present el reglament de la Lliga de Debat Sènior (Xarxa Vives, 2024) per tenir en compte els requisits formals de la participació, l'estructura dels debats i els criteris de valoració del jurat. A instàncies de la Carme, repassàvem sovint el règim de puntuacions que els jutges farien servir, i recordàvem la importància de la presentació de l'exposició amb una tesi sòlida, la qualitat del discurs, la flexibilitat, la demostració del domini del tema, els aspectes formals, la posada en escena i l'actitud de respecte i correcció.

RASTREIG DE LA INFORMACIÓ I CONSTRUCCIÓ DE LA BORSA D'EVIDÈNCIES

Pel que fa al fons de la qüestió, des del primer dia ens vam disposar a cercar, seleccionar i classificar les fonts d'informació, que tant podien ser articles científics de publicacions internacionals com notícies de divulgació científica o tecnològica, textos doctrinals, antecedents filosòfics, opinions periodístiques, notícies de mitjans audiovisuals, imatges fotogràfiques, vídeos o altres. Ja des del principi unes fonts imprescindibles van ser els *films de ciència-ficció* relacionats amb la intel·ligència artificial. El Jaume, expert en cinema, va recollir i comentar alguns d'aquests films més coneguts, que ens serien molt útils com a referències en el muntatge dels arguments.

Les fonts d'informació no només ens havien de ser útils en la preparació dels debats, sinó que haurien de constituir la relació d'evidències que, d'acord amb el Reglament de la Lliga de Debat Sènior 2024 (punt 9.6.3), caldria tenir a disposició del jurat durant el debat, com a prova de les dades invocades. La nostra llista d'evidències final va arribar a ser de quasi dues-centes.

Compilant evidències

«Les fonts d'informació havien de constituir la relació d'evidències que, d'acord amb el Reglament, caldria tenir a disposició del jurat com a prova de les dades invocades.»

SISTEMATITZACIÓ DE LA INFORMACIÓ

Per avançar en la concreció dels diversos arguments del debat i digerir tota la informació, la Carme i la Isabel van proposar un model de *fitxa informativa* per utilitzar-la en la recollida de cada argument. S'hi feia constar quina opinió s'hi defensava (la IA *sí* que és una amenaça, o *no* ho és), la referència o font d'informació, la tesi que es proposava, els arguments de suport i la possible confrontació de l'argument. Es van elaborar 37 fitxes.

EMMAGATZEMATGE ACCESSIBLE

Totes les fonts d'informació trobades van ser recollides i organitzades en carpetes digitals dins la carpeta del núvol Google Drive, accessible per a tots. L'Albert també hi anava recollint tots els documents de treball, com ara les actes de reunions, un glossari de termes d'IA, les fitxes elaborades d'arguments, les aportacions individuals i altres.

CONSTRUCCIÓ DELS ARGUMENTS

El Jaume ens va recordar constantment la importància de preparar una bona argumentació i disposar-ne, independentment de si era per defensar el *sí* com per defensar el *no*. A partir de la seva proposta, vam llegir i comentar un article del diari *Ara* sobre com guanyar un debat fins i tot sense tenir la raó. S'hi recomanava dir coses inesperades, incloure anècdotes personals i, sobretot, mantenir la connexió amb el públic. S'hi aconsellava fer servir la regla del tres: organitzar el discurs de manera que els arguments es desenvolupessin en tres paraules o en tres parts.

LLENGUATGE

Tots els membres de l'equip ens vam mostrar disposats a utilitzar un català acurat en les intervencions. El Jaume ens donava consells freqüents sobre el llenguatge a emprar, i ens convidava a evitar afirmacions absolutes (per exemple «això no passarà mai», «no ho crec», «gens»).

PREPARACIÓ DELS PARLAMENTS

Amb tots els arguments recollits vam començar a preparar els parlaments de cadascuna de les dues posicions, el *sí* i el *no*, respecte de la pregunta de si la IA és una amenaça. Repetidament, comentàvem

que calia ser creatius en la construcció dels arguments, i també que tot plegat havia de ser per gaudir-ne i *ser lúdics* en la preparació dels debats.

Ja a menys d'un mes de la data de la competició, la Isabel i l'Albert van anar preparant successivament les exposicions i refutacions amb els suggeriments i comentaris de tots. Quedava clar que els discursos d'exposició inicial i de conclusions podrien ser força fixos, però que per als torns de refutacions caldria tenir un repertori variat d'arguments que permetés utilitzar-ne uns o altres en funció d'allò que l'equip contrari hagués exposat. Alguns dels suggeriments eren posar-hi més vehemència i citar les pel·lícules de ciència-ficció.

PREPARACIÓ D'IL·LUSTRACIONS GRÀFIQUES

També vam anar concretant la utilització d'algunes imatges en presentacions fetes amb el programa PowerPoint de Microsoft per il·lustrar els parlaments, que haurien de ser molt poques i limitades als torns d'exposició inicial i conclusions. Seria més difícil introduir imatges a les refutacions, les quals dependrien dels parlaments de l'equip contrari.

REPARTICIÓ DELS PAPERS EN EL DEBAT

Dues setmanes abans ja teníem força clar qui intervindria a cada parlament. Per al «sí-amença», la Carme es faria càrrec de l'exposició inicial; la Isabel, de les refutacions i el José Eugenio, de les conclusions. Per al «no-amença», el Jordi assumiria l'exposició inicial i les conclusions, i l'Albert, les refutacions.

ELS ASSAJOS

Assajàvem les intervencions orals en les sessions del grup. En els primers temps, lligides, i més endavant, amb expressió espontània. Ho comentàvem, ho repetíem i ho tornàvem a discutir.

A finals de maig, quan ens va semblar que internament ja ho havíem madurat prou, vam oferir als nostres col·legues de Iubilo URV una sessió pública —una mena d'assaig general o d'entrenament final— a la sala de juntes del campus Catalunya. Hi vam escenificar un enfrontament dialèctic entre nosaltres, tot defensant, uns i altres, les dues posicions possibles respecte del tema proposat. L'acte va ser força estimulant per al col·lectiu en tots els sentits. A banda de felicitar-nos pels parlaments i refutacions, els companys de Iubilo URV ens van suggerir unes quantes millores.

Assaig general

Quan ja ho teníem prou madurat, vam organitzar un assaig general amb públic —els companys de Iubilo— per posar a prova les dues posicions del debat i recollir suggeriments de millora.

ELS SUPORTS

Uns dies abans de la competició, la Carme va preparar i imprimir les *cartolines* amb els guions, les llistes d'*evidències*, l'Albert va preparar les presentacions de PowerPoint amb les *imatges*, i entre tots dos van preparar una llista de *les preguntes* amb les quals, d'acord amb el reglament, interrompríem el discurs dels contraris.

TOT A PUNT!

En els dies previs a la competició vam tancar la feina de preparació amb un assaig final. A aquelles altures, els companys ja ens havíem conegut —o havíem aprofundit en la coneixença mútua—, ens havíem fet amics —o més amics, encara!— i preparàvem la motxilla, els dossiers i les fitxes per anar al Campus Nord de la Universitat Politècnica Catalunya, edifici Vèrtex, on hauríem de representar la URV a la competició sènior de la Lliga de Debat de la Xarxa Vives d'Universitats 2024. El vicerector de Compromís Social i Sostenibilitat de la nostra universitat, Jordi Diloli, ens va desitjar bona feina i bona sort.

3. AMENACES I RISCOS DE LA IA

Dels perills d'un ús inadequat

A més de consumir molta energia, la intel·ligència artificial pot accentuar la desigualtat, degradar l'ètica i generar desconeixement si s'utilitza de manera inadequada.

Tal com hem anunciat, exposarem seguidament els fets i opinions que ens portarien a pensar que la intel·ligència artificial és realment una amenaça per a la humanitat: allò que podríem catalogar com a indicis de perill. En el capítol posterior ens referirem, en canvi, als arguments que descarten que dels riscos anunciats se'n derivi una veritable amenaça. Ho presentem amb els arguments que hem elaborat a partir de les nombroses fonts d'informació trobades, algunes de les quals són assenyalades dins el text, amb hipertext de les referències respectives, i també recollides a la llista de *bibliografia* del capítol 13.

Tot i que intentem fer veure els diferents aspectes de cada idea utilitzada en un sentit o en l'altre del debat, de vegades el llenguatge que emprem és molt directe i pot fer pensar que compartim al 100 % totes les afirmacions que fem. Però cal tenir present que, tant en el cas dels arguments a favor com en contra, les idees exposades aquí no reflecteixen tant la nostra opinió —que, d'altra banda, és plural, ja que no deixem de ser un grup de persones amb diferents pensaments— com la dels autors consultats. I si de vegades expressem alguna idea de forma molt categòrica, pot ser perquè ens hem acostumat a la necessitat de mostrar convenciment en el procés de debat, més que per assumir-la completament.

Tant en aquest capítol com en el següent, comentarem, en primer lloc, les grans idees o arguments que permeten defensar la posició general a favor o en contra, per passar després a exposar aquells plantejaments que permeten refusar els arguments contraris. Com s'ha dit, per preparar el debat era una tasca essencial elaborar aquestes refutacions de la part contrària. El conjunt (idees en favor d'una posició i refutacions dels arguments de la part contrària) és el resultat principal de la feina duta a terme al llarg dels mesos previs al debat i també, en algun cas, introduint reflexions, idees o opinions tretes de les mateixes sessions de debat a la UPC.

Pel que fa específicament a *aquest capítol d'amenaces i riscos de la IA*, ens fixarem en les aplicacions fraudulentament a què dona lloc l'ús actual i previst de la IA, a més de fer veure altres conseqüències negatives que podrien considerar-se, en algun cas, indicis certs de perill i, per tant, donarien peu a considerar efectivament la IA com una amenaça per al conjunt de la humanitat.

La perpetuació de les discriminacions

Qualsevol de nosaltres, en ocasió de fer ús de la IA generativa (conversacional i generació de textos, imatges i vídeos), ha pogut advertir que les bases de dades amb les quals opera estan plenes de prejudicis i d'estereotips de gènere, raça i classe social i que aquests estereotips impregnen el resultat dels textos i les imatges generades. Els assajos que fem amb el ChatGPT (una IA generativa conversacional que, en una versió bàsica, està a l'abast de tothom) ens ofereixen uns resultats que desconeixen la nostra inquietud actual per la igualtat. Si no se li insisteix prou, al ChatGPT li costa mantenir la conversa en llengües minoritàries, i té dificultats per generar la imatge d'una empresària obesa, d'una operària dona, d'un electricista de pell fosca, d'un músic ancià, d'una policia en cadira de rodes. Existeix la por, doncs, que la utilització reiterada d'aquest material esbiaixat perpetuarà les discriminacions que la nostra societat ja pateix avui dia (Hsu, 2024) i suposarà la pèrdua

de la diversitat i de tots els valors que l'acompanyen. Moltes llengües avui encara vives seran fagocitades pels idiomes majoritaris anglès i mandarí. També hi ha el perill que els algorismes de reconeixement facial consolidin les discriminacions que tenim avui dia (Chomsky, 2023).

En la mesura que recull els estereotips de gènere i els difon exponencialment, la IA perpetua els rols associats tradicionalment a les dones i entorpeix la consecució de la igualtat desitjada, tal com comenten Lucia Ortiz de Zárate i Ariana Guevara (2021). Altra-ment, les dones són les víctimes majoritàries de la substitució dràstica de llocs de treball per raó de la implantació de sistemes d'IA (Sempere, 2024).

L'increment del frau en múltiples àmbits

En el camp de la ciència, la IA generativa de textos permet escriure llibres, articles científics, etc. que no són el resultat d'una recerca contrastada, i que permet que els que la utilitzen facin passar com a propis treballs aliens. De resultes d'aquest tipus de frau, els científics han de posar en dubte el material amb què treballen (Nature Editorial, 2024) (figura 1). L'existència d'aquests instruments fraudulents posa en crisi els procediments de competició científica per a l'obtenció de beques, places de recerca, places de docència, etc.

El 2023 l'ús d'IA generativa de text en publicacions científiques s'ha quantificat en un 1 %. La IA genera errors i paràgrafs sense sentit, i és causa de l'empobriment del llenguatge (Degeurin, 2024).



Figura 1. Recreació de la possibilitat que articles científics on s'ha utilitzat la IA siguin fraudulents.

Les trampes que la IA permet fer als estudiants en els exàmens i treballs impedeixen fer-ne una avaluació correcta. Els estudiants, en ignorar quins són els coneixements i les habilitats adquirides, veuen entorpit el seu procés de formació. La docència se n'està ressentint, i la professionalitat dels graduats, també.

En l'art, la facilitat per generar textos (literaris, científics, periodístics, etc.), imatges i vídeos artístics implica menys creativitat en general. Amb la IA les falsificacions, els robatoris i tot tipus de frau aniran a més. La indústria cinematogràfica i de mitjans audiovisuals està en perill. És sabut que ja s'han posat en circulació falsificacions de l'obra d'artistes musicals, com li ha passat a Taylor Swift. També hem vist com a la presentació del nou ChatGPT d'OpenAI es va emetre la veu de Scarlett Johansson en un parlament produït artificialment sense el seu permís.

Hi ha també un perill gran d'influència en l'opinió pública i la política democràtica gràcies a la IA. Ja sabem que la difusió de notícies falses no és un fenomen nou, però les facilitats que la IA ofereix per generar imatges i sons promou la proliferació de les notícies

enganyoses. La creació d'imatges falses a xarxes, filtres de TikTok, etc. dona lloc a tergiversacions de la realitat.

Notícies enganyoses

La creació automatitzada de notícies falses i propaganda representa una amenaça directa per a la transparència i l'estabilitat democràtica.

Altrament, la IA generativa pot induir-nos a pensar que, gràcies a la lectura, en algunes matèries «hi entenem alguna cosa», però el cert és que ja hi ha qui la utilitza per a la difusió d'opinions tendencioses, per difondre propaganda, influir en eleccions i manipular la gent. El potencial de la IA per desinformar, incitar a la violència i condicionar els vots pot destrossar, doncs, les modestes democràcies per les quals es regeixen els interessos col·lectius.

En el camp de l'economia, veus autoritzades afirmen que la IA pot generar frau considerables: per exemple, pot crear falses expectatives en relació amb el valor de les accions en les societats mercantils i pot generar efectes bombolla (Del Castillo, 2024b).

En qualsevol cas, la disposició de màquines intel·ligents ja està fent que molts llocs de treball esdevinguin innecessaris, i tot fa preveure que els acomiadaments deixaran molta gent sense recursos (figura 2).

D'altra banda, la IA pot ser usada per predir el temps restant de vida de les persones. Les companyies d'assegurances se'n poden valer per establir el risc —evidentment en detriment de l'assegurat— en el moment de contractar les pòlisses i establir les primes (Bas-Wohlert, 2024), i hi ha qui alerta de la influència que aquesta predicció tindrà en ocasió de la selecció entre candidats a un lloc de treball, per exemple.



Figura 2. Recreació d'una oficina on només hi ha ordinadors i robots, mentre que a fora la gent fa cua en cerca de feina.

Totes aquestes possibilitats —o certeses— de frau a través de la IA poden modelar comportaments individuals i de grup. [Raul Limón \(2024\)](#) cita Sam Altman, màxim responsable a la companyia OpenAI en aquell moment, i descriu com la IA avançada podria alterar radicalment la naturalesa del treball, l'educació i les activitats creatives, així com la manera com ens comuniquem, coordinem i negociem entre nosaltres, i influir així en «allò que volem ser i arribar a ser». Els agents de la IA poden influir de forma més fàcil que mai en el públic amb la persuasió racional, la manipulació, l'engany, la coerció i l'explotació.

Finalment, hem de citar Noam Chomsky, que ens ha avisat que la IA —sobretot la que es basa en l'aprenentatge automàtic— de-

gradarà l'ètica, ja que incorpora una concepció del llenguatge i del coneixement fonamentalment defectuosa.

Els costos en la salut mental

En una recerca extensa, Amanda [Ruggeri \(2024\)](#) adverteix que la IA provocarà un detriment de la nostra salut mental. En el moment en què la major part dels nostres interlocutors siguin màquines, i quan a l'altra banda del fil del telèfon no hi hagi una persona més o menys empàtica, sinó un mecanisme, la nostra sensació de soledat augmentarà. La IA generativa es pot usar per conversar amb difunts estimats, cosa que pot afectar la salut mental dels usuaris. Els assistents virtuals d'IA generativa conversacional de suport emocional no tenen empatia genuïna, la falsegen i poden portar inclús al suïcidi. Són imminents els robots humanoides que, en vídeo, no diferenciarem dels humans i que es comportaran com si ho fossin ([Gil, 2024](#)). La IA generativa conversacional pot crear falses empaties i pot donar lloc a interpretacions errònies en l'expressió dels sentiments. A l'ésser humà, acostumat a comptar amb la sinceritat de les persones del seu entorn amb qui interactua cada dia, no li serà fàcil adaptar-se mentalment per desconfiar de tota paraula amable.

L'amenaça als recursos de la Terra

Avui en dia el desenvolupament de la IA no és eficient per raó de les dificultats en la fabricació d'unitats de processos gràfics (GPU) i de l'acumulació de poder en les empreses que en monopolitzen la producció. La fabricació dels xips dels macrocomputadors necessita matèries primeres que són escasses i metalls de molt difícil extracció. Els materials de la IA (microxips, semiconductors, etc.) són costosos i majoritàriament no reciclables. La instal·lació d'alguns programes d'IA als telèfons intel·ligents que ja estan preparant algunes marques consumeix una quantitat enorme de memòria RAM ([Gil, 2024](#)).

El funcionament de la IA exigeix grans quantitats de recursos naturals i energètics, i és una amenaça a la sostenibilitat ambiental.

S'observa que l'economia mundial està amenaçada seriosament per la destinació a la IA d'una part molt important dels recursos de la Terra. Les fonts d'energia actuals, inclosa l'energia nuclear, no són suficients per satisfer la demanda energètica que requerirà el ple funcionament de la IA. Només un exemple: el centre computacional que Amazon té als Estats Units consumeix l'energia de les dues centrals nuclears (Del Castillo, 2024a), que s'han construït just al seu costat. També la refrigeració d'aquests centres requereix el consum de grans quantitats d'aigua (figura 3).



Figura 3. Recreació d'una central nuclear produint gran quantitat d'energia per a les empreses d'IA, i consumint molta aigua.

Tot sembla indicar, doncs, que aquesta nova tecnologia incrementa les dificultats per aconseguir la sostenibilitat ambiental i econòmica. Els costos que se'n derivaran i la seva incidència directa en

les activitats de producció, distribució, comerç i consum de béns i serveis —és a dir, en l'economia— seran incommensurables. Potser la factura que passarà a la humanitat no es podrà pagar mai. I no es pot descartar que el col·lapse que temem arribi, també, per la via del fracàs de l'economia; és a dir, per la impossibilitat de satisfer els requeriments de la IA amb els recursos de la Terra.

L'enduriment de la violència

La IA facilita a delinqüents i gent sense escrúpols la suplantació d'identitat amb veu i imatge, i l'accés als béns d'altri ([Ortiz de Zárate & Guevara, 2021](#)). La utilització de la IA en les guerres per dotar l'armament de major potència i precisió mortífera s'ha qualificat d'amenaça certa. Els actuals posseïdors dels sistemes d'IA ja són els més poderosos de la Terra, i el potencial de creixement de la tecnologia n'augmentarà el poder de manera exponencial ([Chomsky, 2023](#)). La IA pot esdevenir molt perillosa, i pot ser una veritable amenaça si s'utilitza a favor dels comportaments individuals, egoistes i violents dels humans i de les organitzacions que es lucren amb enganys i estafes.

Ens preocupa que, recentment, alguns dels millors cervells que treballaven per a la IA hagin abandonat el vaixell per desavinença amb les grans empreses que els havien contractat. El motiu que han expressat és que les empreses han alterat els valors ètics i de controls de seguretat que figuraven en els contractes que els vinculaven, i han dit que ells no es volen veure inclosos en el seu desenvolupament futur. Segons l'excip de Google Eric Schmidt ([Altchek A., 2024](#)), aquests experts no volen haver-se de penedir dels resultats finals de la difusió d'aquesta tecnologia, ni volen que els passi com a Robert Oppenheimer, que es va adonar massa tard de la potencialitat destructiva de la bomba atòmica que ell havia contribuït a crear.

La previsió de superació de la intel·ligència humana

L'amenaça principal detectada és, però, la que es relaciona amb la previsió que la IA arribi a ser la IA general o el que se'n diu la *singularitat*, o sigui, la intel·ligència equivalent a la humana, i fins i tot superior (figura 4). En aquest sentit, diversos experts, com Ray Kurzweil en el seu llibre *The singularity is nearer* (Kurzweil, 2024), afirmen que, amb l'augment exponencial de la quantitat d'informació que gestionen els programes d'IA, cap al 2045 la IA superarà la humana. Adverteix que aleshores la IA actuarà autònomament, i alerta del risc afegit que els algorismes d'autoaprenentatge profund s'encriptaran per seguretat pròpia, de manera que cap humà hi podrà accedir (Estupinyà, 2024). Aquesta previsió, entrevista en el recull d'històries curtes *Jo, robot*, d'Isaac Asimov (1950), o el film *Ex Machina*, d'Alex Garland (2014), configura una situació d'amenaça per a la pervivència de la humanitat.

Hi ha qui vaticina que aquest punt de singularitat s'anticiparà per raó de la introducció dels xips quàntics, que ara s'estan desenvolupant (Abbas, 2024).

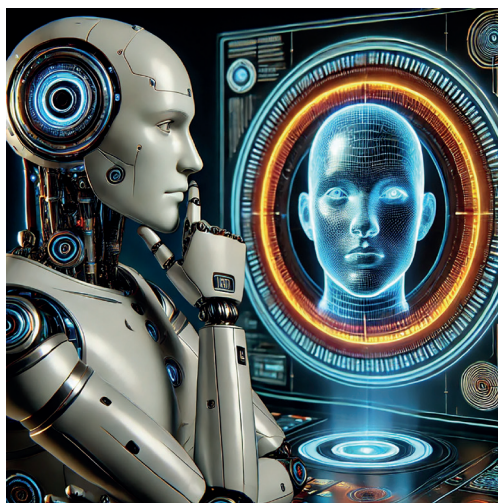


Figura 4. Recreació d'un humanoide d'IA singular en actitud pensativa.

4. OPORTUNITATS I BENEFICIS DE LA IA

Avantatges d'un bon ús

La IA ha arribat per quedar-se i aporta beneficis en múltiples camps, com la salut, la recerca i la productivitat. Els avenços històrics sempre han generat recels, però la capacitat humana d'adaptació i regulació permet transformar els desafiaments en millores socials.

Recull dels principals arguments que descarten l'amenaça

En els capítols anteriors ja s'han exposat algunes idees sobre els aspectes positius de la IA i també sobre aquelles idees i opinions que consideren que la IA no es pot veure com una amenaça per a la humanitat. Toca ara exposar-los de forma resumida i entenedora. D'igual manera com vam fer en el debat, ordenarem aquests fets i idees en tres grans grups, indicarem també alguns punts febles d'aquesta posició, i hi consignarem l'origen de les fonts d'informació, les referències de les quals són a la *bibliografia*, al *capítol 13*.

La IA ja és ben present en l'actualitat

Si la IA es considera una amenaça, vol dir que estem parlant bàsicament del futur. De fet, sovint els arguments en contra de la IA introdueixen, més que fets, possibilitats de futur. O en parlen com si fos un bolet que ha sorgit de sobte fa poc temps. Però resulta que ja fa uns quants anys que convivim amb el que podem incloure, sense cap mena de dubte, dins la intel·ligència artificial. Sense anar més lluny, en la nostra vida quotidiana utilitzem la IA directament o

indirectament quan ens comuniquem (figura 5), fem cerques a internet, traduïm instantàniament des de qualsevol idioma, fem anar programes generatius per escriure i altres aplicacions ofimàtiques que sovint ignorem que són d'IA (Alberca, 2023).

La veritat és que la IA s'ha anat desenvolupant de forma progressiva dins de la informàtica des de fa uns quants anys. Sí que és cert que els darrers dos o tres anys ha tingut un protagonisme molt més gran, gràcies a la difusió de determinades aplicacions com ChatGPT. Però també és veritat que l'abast de la IA com a eina de treball i investigació s'ha difós per múltiples camps, i actualment, doncs, està ben estesa en àmbits com ara l'administració, l'economia, la medicina, la indústria, les comunicacions o, sobretot, en la ciència i la recerca, des de la biomedicina fins a l'astrofísica.



Figura 5. Recreació d'un esquimal en el seu ambient utilitzant el mòbil.

Podem comentar diversos casos que evidencien aquest ús en totes aquestes aplicacions. Per exemple, en relació amb el món de l'empresa industrial i la producció, els avantatges de la utilització de la IA són claríssims i mostren que encara hi ha un llarg camí per recórrer per part de les empreses. Un resum ben clar es pot trobar a [Escala Blog \(2024\)](#). En essència, la IA permet automatitzar i agilitzar tasques l'execució de les quals pot ser avorrida o requerir molt de temps als éssers humans. O sigui, que la IA:

- Permet l'automatització (amb robots controlats per IA) dels processos.
- Potencia les tasques creatives en poder-hi destinar més temps per part dels treballadors.
- Aporta precisió.
- Fa reduir els errors humans.
- Redueix el temps utilitzat en l'anàlisi de dades.
- Facilita el manteniment predictiu de l'equipament industrial.
- Millora la presa de decisions en producció i en comercialització.
- Augmenta la productivitat i la qualitat en la producció.

Tot això ja s'ha començat a aplicar a moltes empreses, tot i que de forma encara desigual. És clar que aquests canvis poden suposar una pèrdua significativa de llocs de treball, com es comenta en parlar de la IA com a amenaça, però recordem que cal fer veure també els llocs de treball que s'estan creant i que es crearan —de major valor afegit— al voltant de la IA.

Els sectors on la IA té més repercussió són el de les tecnologies de la informació i computació, les finances i les vendes, però molts altres també (Andrés, 2024). La majoria d'empreses estan fent cursos accelerats als empleats.

La IA té moltes aplicacions útils —i, per tant, beneficioses—, i se'n van creant de noves cada dia. Un dels camps on aquests beneficis són més clars és la biomedicina, sobretot en la diagnòsi —per

exemple, de diferents tipus de càncer (figura 6)—, i també en l'anàlisi conjunta de diferents proves i dades analítiques dels pacients (Espinoza & Dong, 2020; Huang et al., 2022; NIH, 2024).

I cal reconèixer el paper recent de la IA en les vacunes d'mRNA utilitzades per a la COVID-19 (Dolgin, 2023). A més, cal remarcar que a banda dels coneguts bots de conversa, l'aplicació recent més important de l'aprenentatge profund ha estat AlphaFold 3, el programari de DeepMind (ara filial de Google) per predir interaccions de proteïnes i altres biomolècules (anticossos, ions, RNA, etc.) amb medicaments, fàrmacs, etc. Això ha estat com portar l'alta definició al món biològic, amb una conseqüència evident en la rapidesa de la predicció dels efectes d'un medicament abans d'utilitzar-lo (Abramson et al., 2024). Les farmacèutiques ja estan utilitzant aquesta tecnologia.

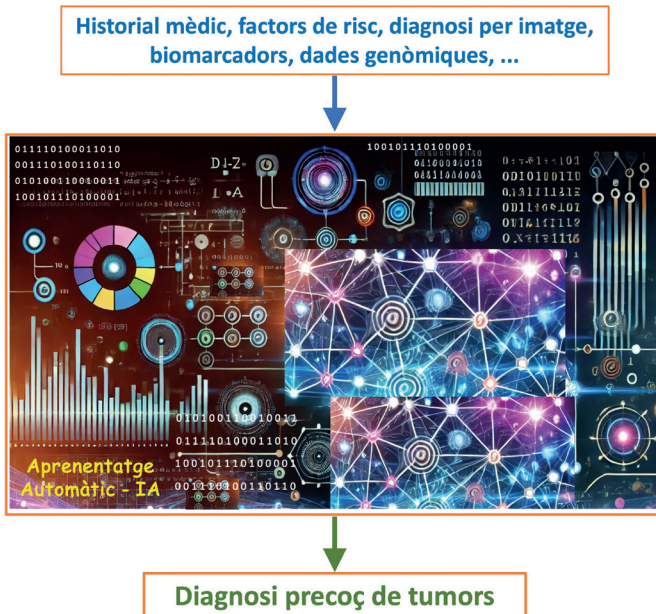


Figura 6. Esquema de l'aplicació d'eines d'aprenentatge automàtic d'IA per a la diagnosi precoç de tumors a partir de múltiples dades del pacient, i per comparació amb les bases de dades clíniques.

Nous horitzons

Els avenços de la IA en salut i ciència, com la diagnosi assistida o la recerca biomèdica accelerada, evidencien el potencial positiu d'aquesta tecnologia.

Però la IA també s'aplica ja actualment en molts altres camps. Vegem-ne alguns exemples de forma breu:

- En astrofísica, programes d'aprenentatge automàtic (*machine learning*) de tractament de *big data* han permès fer un mapa estel·lar de la Via Làctia i passar de 10.000 estrelles a quasi 2.000 milions en quatre anys (Romero-Gómez *et al.*, 2019) (figura 7).
- Molts grups de recerca de diferents àmbits ja estan utilitzant les diverses eines d'IA. Podem veure el recull de les possibilitats a la recerca fet per [Khalifa & Albadawy \(2024\)](#).
- En agricultura, combinant IA amb *big data* obtingudes amb biosensors es pot fer conreu de precisió ajustant l'aplicació de productes químics o l'ús de l'aigua (Borràs, 2024).



Figura 7. Recreació del mapa estel·lar de la Via Làctia.

- En el camp de la història i l'arqueologia hi ha hagut aplicacions diverses. Una de les més notables ha estat el desxiframent a través d'IA de pergamins antics trobats a Herculà, sense necessitat d'exposar-se a perdre bona part de la informació amb la seva manipulació, molt difícil pel fet d'haver sofert el recobriment de cendres de l'erupció i l'efecte tant de la mateixa erupció com del pas del temps posterior (Kuta, 2024).
- Els historiadors de l'art han utilitzat la IA per identificar l'autoria de pintures amb anàlisi visual d'aprenentatge profund, com el fragment de la cara de S. Josep pintat per un altre autor al quadre *La Madonna della rosa* (1520) de Raffaello Sanzio d'Urbino (Nield, 2024) (figura 8).
- No només això. En el camp de l'art, la IA també pot ajudar directament al disseny artístic o als efectes especials en cinema, i contribuir així a la creativitat humana (Lam, 2024). Certament, aquí no deixa d'haver-hi detractors, i es posa en dubte parlar de creativitat quan qui ho fa és una màquina. Però no oblidem que qui crea la màquina és l'ésser humà. El company de la URV Eudald Carbonell comenta precisament que l'ús de la IA creativa ens farà més humans (Carbonell, 2024). En tornarem a parlar.



Figura 8. Quadre *La Madonna della rosa*, de Raffaello Sanzio (Museu del Prado)

En tot cas, el conflicte esclatarà en el moment que, per tal de conjurar els possibles efectes perversos de l'ús de la IA, es proposi prescindir d'algun d'aquests avenços. A què estem disposats a renunciar, si frenem o fins i tot suprimim la IA en determinats àmbits

en nom de la seguretat, la protecció de dades o la desigualtat social? La resposta no és fàcil i segurament requereix una anàlisi complexa, gairebé cas per cas. Però el que queda clar és que si finalment arribem a la conclusió que la IA és una amenaça, ja podem actuar ara mateix per desfer-nos-en, perquè és ben present en molts àmbits de la nostra vida.

La IA no és una amenaça perquè és una tecnologia més

Un dels arguments més estesos entre els defensors de la bondat de la IA és que no és altra cosa que una tecnologia més. Constaten que les reaccions adverses que provoca s'assemblen a les que es produeixen cada cop que hi ha un canvi tecnològic important.

La humanitat ha prosperat a còpia de revolucions tecnològiques. Sense anar més lluny, els darrers 150 anys han estat prolífics en revolucions tecnològiques: electricitat, comunicacions, petroquímica, farmacèutica, biotecnologia, electrònica, internet, els mòbils i ara la IA. Històricament, les aparicions de totes les tecnologies han comportat grans canvis, amb beneficis clars en salut, alimentació, esperança de vida i socialització. En tots els casos hi pot haver alguns efectes no desitjats, però l'experiència mostra, almenys fins ara, que els humans ens sabem adaptar als canvis resultants de les innovacions tecnològiques (Carbonell, 2024).

Aquesta és l'opinió de molts experts, com ara diversos premis Nobel, com Ben Feringa (de Química, el 2016) i Paul Nurse (de Medicina, el 2001) (MacGregor, 2024), experts en computació, com Eric Horvitz (Boyle, 2020) o l'experta Alba Meijide, jove gallega consultora de projectes d'IA (Del Bosque, 2024), que diu el següent:

[...] hi ha molt de fum i de notícies catastrofistes; la fi de la humanitat, la catàstrofe imminent ven molt. En lloc de tenir-ne por, com a societat hauríem de tenir-ne interès, informar-nos-en i forjar una opinió per establir uns límits i uns principis ètics.

Sentit comú vs. Apocalipsi

Les innovacions han generat històricament reaccions de por, però l'experiència demostra que la humanitat ha sabut adaptar-s'hi i aprofitar-ne els beneficis.

En tot cas, la perspectiva històrica fa pensar que, tot i que el canvi que aquesta tecnologia propicia requerirà adaptacions, serà beneficiosa. Els qui descarten que sigui una amenaça acusen els alarmistes de fomentar la por irracional a tot allò que no coneixen. Recorden que això ha passat cada cop que s'ha implantat una tecnologia nova: l'electricitat, internet, els telèfons mòbils...

Com a exemple històric del terror als canvis tecnològics tenim el mite dels «cavalls de ferro», és a dir, la por que van generar els primers trens de vapor (History Skills, 2024). De fet, els primers anys hi va haver algunes explosions de calderes i força accidents. La barreja de classes i gèneres que s'hi propiciava causava escàndol. Corria la brama que, per causa de la velocitat de 50 km/h a què el tren es desplaçava, els passatgers no podrien respirar i que les passatgeres patirien desplaçaments de l'úter. El soroll i el fum a través del camp serien nefastos per a l'agricultura i la ramaderia, i les activitats criminals s'incrementarien a causa de la facilitat de les fugides. Tanmateix, ens podem imaginar la nostra vida actual sense el ferrocarril?



Figura 9. Recreació del terror que podien causar els primers trens de vapor.

La reacció de temor és normal: les situacions desconegudes ens incomoden i instintivament temem que passi el pitjor. En aquest cas, tenim por que la IA porti la humanitat al desastre. Però cal tenir en compte que l'espècie humana ha aconseguit justament els bons nivells actuals de supervivència i l'allargament de l'esperança de vida individual gràcies al desenvolupament continuat i progressiu de noves tecnologies. Per progressar, cal tenir ganes de conèixer i de saber el perquè de les coses. L'atreviment en l'exploració del que és desconegut i la superació de la por al futur ha estat clau en l'expansió i desenvolupament de l'Homo sapiens, com comenta Y. N. Harari en el seu llibre *Sàpiens* (Harari, 2014).

El debat en aquest punt no sembla centrar-se tant en els efectes beneficiosos o perversos de la nova tecnologia, sinó en el possible mal ús que els humans facin de la IA. Com en altres tecnologies, l'aprofitament de la IA pot perjudicar altres persones, sigui amb fraus o altres delictes, o el medi, per sobreexplotació de recursos. Es pot usar com a tecnologia bèl·lica, o per perpetuar estereotips d'ètnia, gènere o classe social. És clar que és possible, com passa amb qualsevol altra tecnologia o estris o eines fetes pels humans. Ha passat sempre. Però en tots els casos parlem d'un mal ús de la tecnologia. L'eina és la responsable de l'amenaça? Per posar un exemple, tot i que es pot utilitzar certament amb finalitats perverses, qualificaríem la informàtica com una amenaça per a la humanitat?

Stephen Hawking, tot i plantejar el risc que una IA sense control podria suposar per a la població i el planeta, reclamava mesures perquè tothom es beneficiés de la IA i perquè la transmissió de coneixement en IA s'equiparés, en importància, a la recerca i a la innovació: cal divulgar els avenços tècnics i científics al públic en general, tant des de les institucions públiques com privades, per tal d'evitar conseqüències no desitjades derivades d'un mal ús (Hawking, 2014).

La IA és artificial, però és de veritat intel·ligent?

El tercer gran argument o grup d'arguments a favor de la innocuïtat de la IA rau en la seva naturalesa. Segons els defensors de la tecnologia, la IA no deixa de ser un conjunt de programes informàtics que simulen alguns aspectes de la intel·ligència humana. De vegades, els alarmistes la presenten com una tecnologia ingovernable.

El cert és, però, que si adjectivem la intel·ligència d'artificial és perquè és una creació humana i que sense els humans no existiria. No deixa de ser intel·ligent, això sí. Però la intel·ligència real, la humana, és molt més que la IA: és el producte del funcionament del cervell humà, que té unes característiques molt complexes que no podem demanar a la IA, almenys tal com la coneixem (Microns Explorer, 2024). La intel·ligència humana inclou el sentit comú, la intuïció, l'autoconsciència, la capacitat d'interactuar socialment, la capacitat d'abstracció, el sentit de l'humor... Podem trobar això en la IA? El cert és que no. Jordi Pompas, expert en IA, no només indica que a la IA li manca creativitat, consciència i emoció, sinó que la seva eficàcia és relativa perquè depèn de la quantitat i la qualitat de les dades que li donem (Pompas, 2023). Per tant, es podria afirmar que la IA no és prou intel·ligent per ser una amenaça.

«Intel·ligència» incompleta

La ment humana compta amb habilitats com el sentit comú, la intuïció, l'autoconsciència, la interacció social, l'abstracció o fins i tot l'humor. De moment, la IA no pot replicar aquestes capacitats.

Ara bé, és evident que la IA no és un sol programa, sinó un conjunt de programes i algorismes cada cop més complexos. N'hi pot haver que siguin absolutament inofensius i, en canvi, n'hi pot haver d'altres que, sigui ara sigui en el futur pròxim, s'assemblin més al cervell humà. Aleshores, on posem la frontera entre les aplicacions

que són innòcues i les que, per raó de la seva proximitat al funcionament del cervell humà, haurien de tenir el creixement limitat per mantenir-les en un nivell de complexitat que en garanteixi el control humà sense problemes?

És important tenir en compte les diferències que existeixen entre les múltiples aplicacions de la IA. Ara, quan parlem d'IA només se'ns acudeixen les aplicacions d'IA generativa conversacionals o bots de conversa —com el ChatGPT—, que simulen aspectes de llenguatge. Però, per exemple, les calculadores són estris artificials, inventats pels humans, que simulen i reproduïxen aspectes de lògica i matemàtiques de la intel·ligència humana i, per tant, podríem qualificar-los d'IA, no? En tenim des dels antics àbacs a les calculadores científiques, i ara comptem amb supercomputadors com el Mare Nostrum 5. D'altra banda, una gran quantitat d'IA d'avantguarda s'ha filtrat en aplicacions generals (ex. cerques de Google o etiquetatge de cares a Facebook), sovint sense que s'anomeni IA perquè una vegada que alguna cosa es fa prou útil i prou comuna, ja no s'etiqueta com a IA (Kaplan *et al.*, 2019).

Tanmateix, ara com ara, el fet que la IA pugui arribar a ser autònoma i independent dels humans —com s'alerta en determinades opinions o, més sovint, en pel·lícules de ciència ficció— no sembla gaire pròxim. Per exemple, la IA dels bots de conversa actuals no és gaire espavilada. Pot ser gaire intel·ligent si és incapaç d'entendre els acudits de Groucho Marx? (Schank, 2017).

A més del sentit de l'humor, com hem dit abans, hi ha força capacitats de la ment humana que les màquines amb IA no tenen. De fet, la nostra comprensió del funcionament del cervell humà encara és a les beceroles. Però en tot cas, queda clar que la manera com els sistemes d'IA aprenen és més simple que no pas la manera com aprèn el cervell humà. Emular tot el seu funcionament amb IA, o sigui, la creació d'una IA forta o general que actuï exactament igual que els humans és, ara per ara, hipotètic. No hi ha cap programa d'IA que es pugui considerar intel·ligent en termes generals. La suma d'un munt d'IA estretes o febles (aplicades a programes

específics per a una tasca definida) mai serà una IA forta o general (News and Events, 2024).

Les IA més avançades, com per exemple el reconeixement d'imatges, són de xarxes neuronals artificials convolutives d'auto-aprenentatge profund multicapes, però d'aprenentatge supervisat, no comparable ni al dels infants (Mitchell, 2024). Les màquines d'IA canvien gradualment els encerts a mesura que processen els exemples del conjunt de dades d'entrenament un cop i un altre, al llarg de moltes repeticions amb les dades d'entrenament, i aprenen a classificar cada entrada dins d'un conjunt fix de possibles categories de sortida. Els infants, al contrari, ja de ben petits aprenen un conjunt obert de categories i poden reconèixer casos de la majoria de categories després de veure'n només alguns exemples. A més, no aprenen de forma passiva: fan preguntes, demanen informació sobre tot el que els desperta curiositat, dedueixen abstraccions i connexions entre conceptes i sobretot exploren el món.

Tot i tenint en compte les reserves expressades sobre la capacitat de la IA, el seu desenvolupament sembla tan ràpid que hi ha qui pensa que serà difícil reaccionar al moment en què això passi. És cert que aquests darrers anys els progressos han estat ràpids, però, vista la complexitat del cervell humà, tot plegat no sembla que hagi de ser tan veloç i, com en altres ocasions, probablement els humans tindrem temps d'adaptar-nos-hi i en traurem profit, com ha passat en altres revolucions tecnològiques.

Resumint, els qui creuen que la IA no és una amenaça per a la humanitat són els qui sostenen que aquesta tecnologia servirà per millorar les habilitats humanes i la convivència, per avançar en la recerca, per satisfer les necessitats de salut, alimentació i educació, i millorar, per tant, la qualitat de vida en general (Lawlor & Chang, 2023; [Sánchez Zaplana](#), 2023).

5. REFUTACIONS PRINCIPALS ALS ARGUMENTS A FAVOR I EN CONTRA DE LA IA

Preparant les rèpliques

Per debatre amb solidesa cal anticipar els arguments rivals. Aquest capítol recull les refutacions clau preparades per qüestionar tant les visions optimistes com les alarmistes sobre la IA i demostra com la dialèctica pot matisar el debat i posar en evidència els punts febles de cada posició.

L'equip havia de preparar una borsa de contraarguments en relació amb cadascun dels posicionaments sobre la IA que eren previsibles, a fi de poder-los esgrimir oportunament davant dels adversaris. La lectura curosa de les fonts que prèviament havíem seleccionat ens va proporcionar el substrat que ens havia de permetre improvisar les refutacions als arguments de l'equip contrari. De les que vam preparar i utilitzar en el debat exposem, en aquest capítol, les que considerem més importants. Fan referència, lògicament, a qüestions que ja s'han comentat, però introdueixen idees, dades i punts de vista diferents que enriqueixen el debat en un sentit o en un altre. Compte, que cal llegir-los com el que són: contraarguments a utilitzar en un debat sobre la IA, tal com els vam crear. Per tant, les afirmacions o negacions que s'hi fan no són gens matisades. Hem cregut que valia la pena deixar-ho així —en el llenguatge del debat— per entendre com es prepara i es planteja la defensa d'un posicionament o altre en una competició com aquesta.

Refutacions als arguments dels defensors de la IA

En aquest primer gran apartat del capítol aprofundirem en aquells aspectes que permeten rebaixar la versemblança dels arguments esgrimits pels defensors de la IA com un element positiu per a la humanitat i no com una amenaça.

MANCA D'EINES DE CONTROL GLOBAL

Els defensors del posicionament que la IA no és cap amenaça solen sostenir que la humanitat és capaç de controlar els efectes perversos de la IA.

Si el govern de tota la Terra estigués en mans d'una única organització democràtica que vetllés pel bé comú de tots els éssers, potser podríem confiar en l'encert d'aquesta entitat per controlar els efectes nefastos de la IA. Això, ara per ara, no és factible, vist el poc cas que els estats fan de les resolucions de l'Organització de les Nacions Unides. De moment, no hi ha una organització única i democràtica tutora dels interessos de la humanitat sencera. Fins i tot els tractats entre nacions són massa febles per obligar els tenidors dels recursos d'IA que els apliquin únicament i exclusivament a aconseguir el benestar global. Hi ha qui considera insuficient el reglament de què s'ha dotat la Unió Europea, per exemple, ja que no té altra ambició que preservar la lliure competència del nostre mercat en el comerç de recursos d'IA.

La IA, capaç de fer càlculs i processos a una velocitat que fins ara no era imaginable, podria multiplicar per molt la capacitat de manipulació que els poderosos ja tenen avui en dia. La manca d'una organització mundial que preservi els valors de la humanitat en deixa la gestió en mans privades, cosa que ens obliga a plantejar-nos quin tipus de persones o entitats són els amos dels enginyers, patents i saber fer d'aquesta tecnologia. Tot fa témer que, mentre les corporacions dotades d'IA colpegin a tort i a dret les nostres civilitzacions i es carreguin els patrimonis culturals aconseguits per

totes elles durant mil·lennis, el poder de l'oligarquia serà cada vegada més gran, i l'explotació de l'ésser humà per l'ésser humà serà finalment irreversible.

Si els defensors del posicionament que la IA no és una amenaça argumentessin que l'expansió de la IA pot ser molt beneficiosa per a la humanitat, els recordaríem que la IA és una eina de manipulació molt eficaç. Sam Altman, màxim responsable d'OpenAI en aquell moment ([Limón](#), 2024) identifica diverses vies de manipulació. Una primera via és la de la manipulació de l'usuari per part de l'agent de la IA, en funció dels objectius de la companyia propietària, mitjançant informacions falses, per exemple. Una segona via: que l'usuari utilitzi l'agent d'IA per generar una posició de domini. I una tercera via: que sigui una entitat qui manipuli l'usuari per tal de satisfer un suposat interès col·lectiu.

TAL VEGADA ENS DIRAN QUE LA IA ÉS UNA TECNOLOGIA MÉS

Per a nosaltres, la IA no és, simplement, «una tecnologia més». Els bots d'IA generativa conversacional de suport emocional no tenen empatia genuïna, la falsegen i poden portar al suïcidi ([Ruggeri](#), 2024).

La darrera cimera de Davos ha posat de manifest que la IA està substituint cada vegada més treballadors i no es descarta que en un futur proper no ens deixi a tots sense feina. N'alerten els principals experts en IA, com Mustafa Suleyman, cofundador de DeepMind, entre d'altres ([Daniel](#), 2024).

La IA ja està sent utilitzada amb intencions destructives. Armes autònomes acoblades a drons que tenen programats els objectius amb IA donen lloc a un augment de violència i d'inseguretat ([Israel](#), 2023). És altament previsible el desencadenament de conflictes geopolítics per raó de les matèries primeres necessàries per als semiconductors. No es poden descartar guerres Xina-EUA-Rússia, amb punt neuràlgic a l'Orient Mitjà.

No es pot descartar que l'expansió de la IA posi en crisi l'objectiu de benestar que persegueix la humanitat. És, doncs, efectivament, una amenaça.

LA PRESENCIA EFECTIVA DE LA IA EN LA NOSTRA SOCIETAT

Algú ens pot objectar que actualment ja fem la IA en molts camps —l'administració pública, l'economia, la indústria, la ciència, la recerca...— i que ha passat a ser una eina quasi rutinària.

Cert és que la IA s'utilitza habitualment en la medicina, la seguretat, l'administració i molts altres camps, però aquesta rutina no la fa necessàriament saludable. Qui no ha patit «ansietat algorítmica», en els termes que Kyle Chayca ([Pérez-Colomé, 2024](#)) defineix aquesta alteració, quan ha de fer un tràmit quotidià en línia? Quan ha de formular, per exemple, una reclamació a una entitat bancària? Certament, la interlocució amb les màquines —que són eines privades de bagatge emocional— pot repercutir negativament en la salut mental de les persones. La humanització de la màquina pot crear falses relacions de confiança i vinculació emocional. Citem el cineasta Spielberg, que recull aquest fenomen al film *Intel·ligència artificial*, estrenat ja fa més de 20 anys!

LIMITACIÓ DELS AVANTATGES ATRIBUÏTS A LA IA

Si els avantatges de la IA (que, certament, n'hi ha) estiguessin a l'abast de tothom, potser sí que seria una tecnologia sostenible. En tal cas, potser sabríem acudir a l'aplicació del sentit comú —és a dir, podríem confiar en la sensatesa compartida que tots els humans tenim—, per posar un límit natural a l'expansió perillosa de la nova tecnologia. Però no hem d'oblidar que són només unes poques persones —unes poques empreses— les que exploten la IA en règim d'oligopoli. Ens ho adverteix l'expert Michael [Wooldridge \(2024\)](#). La IA pot incrementar el mal comportament de la humanitat, de la mateixa manera que ho va fer la bomba atòmica d'Oppenheimer.

Hem de recordar que, des de molt antic, escriptors com Plaute i pensadors com Erasme i Hobbes han advertit que l'home és un llop per a l'home: *Homo homini lupus*. Hem d'estar alerta. No hem de deixar que les nostres identitats individuals resultin suplantades per sistemes d'IA, ni hem de tolerar que les nostres dades siguin tractades pels nostres patrons per mesurar la nostra força de treball ni per les asseguradores per predir el temps que conservarem la vida (Ruggeri, 2024). Si insistim que la IA *sí* que és una amenaça per a la humanitat és perquè, encara que presenta alguns avantatges, no és previsible que siguin extensius per a tothom.

PER CONCLOURE, TORNEM AL PUNT DE PARTIDA

Si acudim a la filosofia, veiem que la finalitat i el sentit de la humanitat no és de créixer exponencialment en nombre d'individus. Ni tan sols ho és d'augmentar els anys de vida de cadascun d'ells. L'objectiu de la humanitat és, simplement, d'aconseguir el bé comú: que el major nombre possible de gent visqui en condicions de benestar. Amb Aristòtil constatem que la condició plena de l'ésser humà i de la polis només es pot adquirir a través de la col·lectivitat: la comunicació amb els altres, la col·laboració i la compartició de coneixement.

No creiem que la IA ens hagi de ser de gaire ajuda en la consecució d'aquesta fita, ja que no és capaç de tenir perspectives personals, no actua en funció de creences i no pot comprendre exactament el sentiment que hi ha darrere d'un text. Perpetua els biaixos que ja existeixen en matèria de gènere, races i diversitats. Degrada la nostra ètica.

L'autonomia de la IA és un mite: no té consciència ni valors propis, però el seu ús irresponsable pot fer que els objectius de la humanitat fracassin.

La IA presagia que el nostre procés d'adaptació a l'entorn deixarà de dependre del nostre intel·lecte natural —sempre enginyós, sempre a punt per adaptar-se, sempre amatent a la innovació— i quedarà limitat a les possibilitats que ofereixi un sistema artificial que, per ampli que sigui, estarà simplement nodrit d'experiències velles farcides de prejudicis. Sense referències en relació amb què és cert i què no ho és, la humanitat perd una eina bàsica per a la gestió dels seus interessos.

El dia que el resultat de la intercomunicació, la col·laboració i la compartició del coneixement en el context de la *polis* aristotèlica —és a dir, en el si de la societat formada per persones— no trobi el camí per expressar-se, la humanitat, certament, sucumbirà.

Refutacions als arguments dels detractors de la IA

Al llarg de l'exposició d'arguments feta en els capítols anteriors ja han anat sortint, inevitablement, algunes de les refutacions principals als arguments dels detractors de la IA, els que la veuen com una amenaça. Malgrat això, aquí incloem de forma breu algunes de les més destacades. Com en l'apartat anterior, exposem les refutacions tal com van ser pensades per al debat, tot i que no les vam haver d'utilitzar perquè no ens va tocar de defensar aquest posicionament; les afirmacions i negacions són, doncs, també matisables, però hem preferit de mostrar-les tal com les havíem preparat.

SOBRE LA IA COM A PERPETUADORA DE DISCRIMINACIONS

Als qui veuen com una amenaça per a la humanitat la perpetuació de les discriminacions que derivaria del fet que les bases de dades emprades per la IA siguin plenes de prejudicis de gènere, raça i classe social, se'ls ha de dir que la IA reproduceix el món tal com funciona ara. És clar que cal introduir correccions en el seu funcionament per evitar aquests biaixos. Respecte al biaix de gènere, la dona ja es va incorporant a tots els àmbits en la majoria de països i, de fet, el

talent natural universal —masculí i femení— proporcionarà antidots contra qualsevol excés que la IA pugui cometre.

Quant al perill de fagocitació dels idiomes minoritaris —i/o minoritzats, com el català— per part dels majoritaris (anglès, castellà, mandarí, etc.), aquest no és un fenomen provocat per la IA, sinó que és conseqüència de la inevitable globalització. En canvi, les eines d'IA disponibles ja avui en dia permeten justament mantenir les llengües minoritàries en qualsevol ocasió, escrita o oral, independentment del receptor o lector que no conegui la llengua i viceversa. En el cas del català, cal recordar que malgrat que es troba en el 75è lloc mundial per parlants, és la 35a llengua en presència al món digital ([W3Techs, 2024](#)), i això ha anat augmentant aquests darrers anys gràcies a les eines d'IA. Per tant, aquest és un avantatge que pot ajudar al manteniment i vitalitat de tots els idiomes i, en particular, del nostre.

SOBRE LA IA COM A FACILITADORA DE FRAUS I ENGANYS

Als qui pensen que la IA és una amenaça perquè facilita el frau en molts àmbits, incloent-hi suplantacions d'identitat i altres danys, també en el camp científic, i enganyos en l'economia, tot fent que molts llocs de treball esdevinguin innecessaris, se'ls respon que una de les àrees on més s'està desenvolupant la IA és, justament, la de ciberseguretat. Aquesta disciplina, que inclou la criptografia, és perfectament adequada per a la detecció i correcció de fraus, fins i tot en l'àmbit policial, tal com detalla l'expert José Luís Ponce ([Ponce, 2024](#)). La IA s'aplica ja per a la detecció i identificació ràpida d'amenaques, la gestió de contrasenyes, l'anàlisi forense, l'autenticació biomètrica, la prevenció dels nous ciberdelictes i d'altres. També s'utilitzen programes d'IA per detectar les notícies o imatges falses (Berrondo-Otermin [et al.](#), 2023).

Respecte als possibles fraus en la ciència, pensem que ja està havent-hi una regulació de l'ús de la IA. Per un costat, els mateixos investigadors i professionals són conscients d'aquests riscos i de la

conveniència de contrastar sempre qualsevol dada, com en tots els àmbits, i per això s'utilitzen cada cop més programes detectors de possibles frauds. En aquest sentit, moltes revistes científiques estan introduint sistemes de control dels usos, sobretot de la IA generativa (Kacena et al., 2024) i, per exemple, l'empresa editora Wiley ha tancat 19 revistes i ha retirat 11.000 articles per mal ús de la IA (Stéphane, 2024).

Contra el frau

Si bé es pot fer servir per enganyar, la IA també és eficaç en la lluita contra els ciberdelictes, les notícies falses i les suplantacions d'identitat.

Respecte a l'argument que molts llocs de treball seran innecessaris i que hi haurà molts acomiadaments, cal dir que la IA és una gran revolució tecnològica i que, com en les revolucions anteriors, va lligada a una renovació del perfil dels treballadors que es necessiten. S'eliminen llocs de treball i se'n generen molts d'altres. A banda, cal recordar que els acomiadaments els fan els humans, no els fa la IA. De fet, ja hi ha demanda de feines noves relacionades amb la mateixa IA (Patrizio, 2024). Són llocs de treball que sovint requereixen una bona formació en informàtica, però també n'hi ha alguns on es demanen coneixements de psicologia, de seguretat, de privacitat i d'ètica. Els sectors principals on s'estan creant llocs de treball nous són la tecnologia de comunicacions, les finances, la recerca en salut, el comerç i la producció industrial. Els avantatges de la utilització de la IA a l'entorn empresarial són claríssims (Escala Blog, 2024), i la majoria de grans empreses ofereixen cursos accelerats als empleats, ja que la IA permet automatitzar i agilitzar tasques l'execució de les quals podria ser monòtona i avorrida, o podria requerir molt de temps als humans. La nova tecnologia augmenta la productivitat i la qualitat en els productes. En acabat, com recorda l'expert en visió computeritzada Santiago Valdarrama: «La IA no et substituirà, ho farà una persona que utilitzi la IA» (Rufus, 2023).

SOBRE ALTRES EFECTES NOCIUS DE LA IA

Altres arguments que s'esgrimeixen en favor de la tesi que la IA és una amenaça són els dels costos en la salut mental, l'economia i l'enduriment de la violència.

Com a contrapunt dels evidents beneficis comentats de la IA, totes les noves tecnologies tenen alguns costos o aspectes negatius que cal minimitzar. A banda de la necessitat de regulació i control, que ja s'està establint a molts àmbits, tant de les administracions com dels mateixos productors i consumidors per les possibles afectacions en la salut mental, és clar que cal educació i divulgació perquè tothom sigui conscient dels usos correctes d'aquesta tecnologia. En aquest sentit, hi ha les guies de la [Unesco \(2024\)](#) o la recent Guia 2024 de la UOC ([Angulo, 2024](#)).

Respecte als costos econòmics de la IA, si bé és veritat que ara s'estan emprant molts recursos, tant energètics com materials, per al funcionament dels grans supercomputadors que gestionen la immensa quantitat de dades utilitzades per tots els usuaris dels programes d'IA, hem de tenir en compte que la seva eficiència també augmenta exponencialment amb la incorporació de sistemes d'estalvi de circuits i de millora de rendiment. De fet, l'evolució positiva del mercat i les fortes inversions en R+D que s'estan duent a terme fan que els microprocessadors siguin cada cop menys costosos. A més, ben aviat estarà disponible la nova generació dels xips quàntics, que permetran accelerar exponencialment les capacitats computacionals i abaratir els costos ([Abbas, 2024](#)). Finalment, els recursos econòmics que s'hi empraran no seran tants ni hi haurà perill d'enfonsament de l'economia mundial, ans al contrari, ja que la IA l'afavorirà, com a revolució tecnològica que és.

Respecte al suposat enduriment de la violència, és evident que la IA és un conjunt d'aplicacions informàtiques que acceleren molt la realització de tasques o faciliten l'ús quasi autònom de moltes màquines, i que algunes d'aquestes màquines poden ser armes. Però, com passa amb les armes ja existents, són eines dirigides per

alguna persona. Per tant, som els humans qui som més o menys violents, sigui individualment o col·lectivament. Com passa en relació amb qualsevol eina, el seu ús ben controlat no necessàriament ha de donar lloc a un enduriment de la violència ja existent.

SOBRE LA CAPACITAT DE LA IA PER SUPERAR LA MENT HUMANA

Finalment, cal afrontar l'argument primordial dels que pensen que la IA és una gran amenaça per a la humanitat perquè es preveu que en poc temps esdevindrà la IA general, o sigui, l'equivalent a la humana i fins i tot superior, segons ho preveuen alguns experts com [Kurzweil \(2005\)](#). Evidentment, és una suposició alarmista basada en la por del que és desconegut i en les pel·lícules de ciència-ficció. Cal tenir en compte, però, que encara no entenem la majoria dels mecanismes neurofisiològics que donen lloc a la intel·ligència humana. La simulació del cervell humà —tan complex!— està molt lluny d'aconseguir-se amb la IA. Es poden arribar a simular força aspectes de la intel·ligència humana, però no la totalitat. Com diu l'empresari informàtic [Mitchell Kapor \(Andersen, 2014\)](#), la intel·ligència humana és un fenomen meravellós, subtil i mal conegut, i, per tant, no hi ha perill que es dupliqui a curt ni a mitjà termini. El mateix Kurzweil també preveia que cap al 2020 els nanorobots ja haurien escanejat tot el cervell i n'haurien entès el funcionament, i suposaven que, per enginyeria inversa, tot ell es podria reproduir amb IA. Ja han passat quatre anys de la data anunciada i no ha passat res d'això. En qualsevol cas, la IA és un conjunt d'eines dissenyades pels humans, i som nosaltres els qui controlem què fa o què pot fer la IA. Per tant, davant de qualsevol sospita d'autonomia o independència de la IA respecte dels humans, o de la detecció d'una amenaça real, sempre hi haurà la possibilitat de desconnectar-la.

Hi ha, finalment, una pregunta que ens fem tant els defensors del *sí* com els defensors del *no*: *la IA és, de veritat, intel·ligent?*

Els experts en la matèria no es posen d'acord. Ara bé, després de la molta informació que hem recollit per realitzar el debat de la

Xarxa Vives —i que dona peu a aquesta publicació—, la reflexió dominant és que la IA no pot modular les seves conclusions. Estem parlant d'un sistema que és capaç de generar informació —molta informació, certament—, però que no pot emetre opinions ponderades, ni està positivament condicionat per creences ni valors. La IA no és conscient de si mateixa ni té perspectives personals (Chomsky, 2023). No pot comprendre els motius d'una decisió. En aquestes condicions, seria desencertat confiar en eventuais decisions morals de la IA generativa. I la resposta a una d'aquestes qüestions morals és, precisament, la que constitueix el tema d'aquest debat: és la IA una amenaça per a la humanitat? És a dir, perdrà la humanitat els seus trets idiosincràtics si en comptes de regir-se pel progrés continuat i col·lectiu de l'enginy natural dels seus components es regeix per una combinació matemàtica?

6. CONCLUSIÓ: LA IA ÉS UNA AMENAÇA O NO?

L'amença és el mal ús

La IA pot ser un risc si reforça desigualtats, facilita frauds o escapa del control humà. Però també és una eina útil en molts àmbits. L'impacte final dependrà de l'ètica i els valors amb què la utilitzem.

Com hem vist, en la preparació de la Lliga de Debat sobre aquest tema vam trobar prou arguments tant en un sentit com en l'altre, i teníem ben preparats els dos posicionaments del debat, malgrat que a la competició només vam tenir l'oportunitat de defensar que la IA és una amenaça, per la mala sort en el sorteig dels tres enfrontaments que vam haver de fer.

Com a conclusions generals sobre si la IA és una amenaça o no, podem recordar breument els arguments principals de les dues posicions.

La IA és una amenaça perquè:

- Genera textos o imatges a partir de bases de dades que mantenen les discriminacions i reforcen els estereotips i els biaixos de gènere, raça i classe social.
- Facilita els frauds en molts àmbits, com les suplantacions d'identitats, les comunicacions, la ciència, l'educació, l'art, l'opinió pública, la política i l'economia.
- Té uns costos mai vistos de funcionament de recursos i energia, en salut mental, en l'economia i en l'enduriment de la violència.

- Pot arribar en pocs anys a ser la IA general o singular, equivalent a la intel·ligència humana i fins i tot superior, amb conseqüències de domini sobre els humans.

La IA no és una amenaça perquè:

- Ja és ben present en l'actualitat i la fem anar quasi sense adonar-nos-en, perquè ens és molt útil en molts camps, començant per la vida quotidiana i en les comunicacions, fins a qualsevol àmbit econòmic, tecnològic i científic.
- És una tecnologia més, que suposa una revolució que comporta grans canvis, la majoria dels quals produeixen beneficis en l'economia, la salut, l'alimentació, l'esperança de vida i la socialització.
- La IA no és tan intel·ligent com els humans, ni ho serà, perquè és un conjunt d'aplicacions en base informàtica i matemàtica que són dissenyades pels humans i, per tant, aquests en mantenen el control.

Amb això, podem dir que el fet que la IA sigui o no una amenaça per a la humanitat dependrà de dos factors:

- Haver acordat que l'objectiu principal de l'existència humana sobre el planeta és el «bé comú» i els valors associats.
- Admetre que això comporta un compromís ètic de l'ésser humà en la fabricació i utilització de les eines.

En definitiva, la IA, com qualsevol instrument, no és per si mateixa una amenaça ni deixa de ser-ho. Els instruments o artefactes no tenen categoria moral, és a dir, no són en ells mateixos ni bons ni dolents; és *l'ésser humà* qui els pot fer servir tant per fer el bé com per fer el mal.

Per tant, l'amenaça prové de la humanitat mateixa, de com utilitzi l'ésser humà la IA i amb quina finalitat. En cas que ho fes per dominar o sotmetre els seus semblants, per trastornar la seva convivència o autodestruir-se en algun aspecte com a ésser humà, llavors sí que seria una amenaça i caldria protegir-se'n.

Regular i educar

No es tracta de prohibir la IA, sinó d'establir mecanismes reguladors i de control efectius, així com d'educar la població sobre els seus beneficis i riscos.

I és que ningú dubta que la capacitat d'actuació de la IA en camps potencialment perillosos per a les persones o el planeta és gran. La IA pot ajudar a perpetuar un sistema d'explotació de recursos absolutament insostenible com el que funciona actualment. O pot aplicar-se en armes cada cop més sofisticades per destruir territoris o exèrcits, com sembla que s'està fent ja en conflictes actuals. És que la solució per evitar els efectes negatius de la IA passa per no desenvolupar-la, i fins i tot prohibir-la? Si fem un repàs a les revolucions tecnològiques anteriors, sembla que la solució hauria de dependre més aviat de dues mesures clares:

- 1) la *regulació* del seu desenvolupament amb mesures efectives de control i,
- 2) el més important, *educar* convenientment la població sobre les possibilitats i els riscos de la IA.

Així doncs, cal confiar en el futur d'un bon ús de la IA, que ajudi a preservar la natura i, per tant, la humanitat (figura 10).



Figura 10. «Futur naturoide»: una imatge d'esperança

Finalment, i com a curiositat, podeu veure al *capítol 12* unes preguntes generals que l'Albert ha fet a l'aplicació d'IA generativa conversacional ChatGPT (versió lliure v2, subscripció via Google):

- Fes un text d'unes 300 paraules que respongui a la pregunta següent: La intel·ligència artificial és una amenaça per a la humanitat?
- Fes un text d'unes 200 paraules màxim que comentis els principals avantatges de la intel·ligència artificial per als majors de 65 anys.

Com podeu veure, la primera resposta compleix força bé l'argumentari que hem fet anar als nostres debats respecte a les dues posicions. I a la segona pregunta que li hem fet, donada la nostra condició de personal jubilat, així com la resta del companys de Iubilo URV i d'altres amics, ChatGPT ens respon amb unes idees força interessants.

7. CONCLUSIONS DE LA NOSTRA EXPERIÈNCIA

El joc com a aprenentatge

A la Lliga de Debat, el sorteig va marcar el posicionament de l'equip i va evidenciar com l'atzar influeix en la vida. Però el joc no és només diversió: és una eina de coneixement, debat i creixement intel·lectual recomanable.

Els missatges amb què els nostres majors ens van criar deien: «Nen, aplica't si vols ser un home de profit»; «Nena, lleva't d'hora i no perdís el temps si vols tenir èxit en la vida». Condicionats per aquestes cantarelles, quan fem el resum de les fites aconseguides en la vida professional ens descuidem sovint de consignar-hi la influència que la sort ha tingut en els nostres èxits. Rara vegada ho hem fet constar en els currículums que hem hagut de confegir durant la nostra vida laboral i acadèmica.

Però ara, ja grans, amb la lucidesa de la senectut, ens adonem que, en els nostres destins, la casualitat ha tingut més pes que no pas la vida esforçada. La nostra sort no només és d'haver nascut en el segle xx en els bonics paratges de l'arc mediterrani, en famílies menestrals amants de la cultura, sinó, sobretot, d'haver pogut triar innumerables vegades entre els diferents camins que la fortuna ens presentava. En el caos de l'Univers som el fruit de la casualitat. Som on som —i som com som— per raó dels milers d'opcions que hem anat adoptant en la vida, més aviat guiats per la intuïció que no pas per raonaments lògics. Opcions que ens han fet descartar les miríades de possibilitats que se'ns han anat presentant i que ens han fet tancar moltes portes que, si les haguéssim obert, ens haurien conduït a resultats vitals diferents dels que tenim.

Ni les lleis de la natura ni la voluntat humana regeixen la fortuna. Però l'atzar no només és cosa de la bruixa: les ciències hi comp-ten, i se'n serveixen. Les matemàtiques, l'estadística, els algorismes, el sentit de la computabilitat de Turing, la teoria de Solomonoff, la mecànica quàntica, l'evolució química, astronòmica, geològica i biològica, la selecció natural, la deriva genètica, etc. Ei! Com s'ho faria el ChatGPT sense l'atzar?

En la nostra participació en la Lliga de Debat Sènior 2024 que organitza la Xarxa Vives d'Universitats, la casualitat ha tingut un paper protagonista: l'atzar ha volgut que el 100 per 100 dels sortejos dels quals depenia el nostre posicionament en la competició hagin tingut un resultat idèntic: sempre ens ha tocat defensar que la intel·ligència artificial sí que és una amenaça per a la humanitat. De res han valgut els afanys de l'equip en la preparació d'arguments per al *no*. La intensitat del treball que hi hem dedicat no ha tingut cap recompensa. La moral de l'esforç ha mostrat la seva opacitat enfront de l'esplendor de Fortuna, la deessa que ha condicionat els resultats dels sortejos d'abans de cada partit.

Per a la humanitat, el fat és un misteri molt atractiu, i la presència d'aquest factor en el joc fa que l'activitat lúdica hagi estat, des de sempre, alguna cosa més que un passatemps. A més, la psicologia considera que el joc és un factor bàsic en el desenvolupament humà, ja que promou aprenentatges i permet adquirir destreses, coneixements i habilitats. El joc és clarament una eina de coneixement.

Ha estat en un context de joc que, aquest 2024, ens hem afrontat a una disciplina nova de trinca en la qual fins ara érem absolutament llecs: la intel·ligència artificial. Ens hi hem aproximat amb esperit lúdic i des de la ignorància. En qualsevol altre registre no hauríem gosat pontificar sobre les amenaces que aquesta nova tecnologia pot arribar a suposar per a la humanitat. Hem fet com els xiquets: per la via del joc ens hem apropiat de broma a un fenomen poc conegut i, jugant, jugant, ens l'hem fet una mica nostre.

Som vius

«I ens direu: Què feu, tan ganàpies, que jugueu com criatures? Us respondrem: Sí, som grans, però som vius! I, com que vivim, juguem!».

El joc, per altra banda, no només és una activitat recreativa — que també ho és—sinó que és altament creativa. El format et permet afrontar un tema seriós amb una actitud amistosa i desenfadada. Per participar-hi has d'obrir la ment i has de mirar d'avaluar i assumir tots els posicionaments possibles. Has de descartar les idees prefixades i els fonamentalismes. T'has de carregar d'evidències científiques per sostenir tots i cadascun dels arguments que preparis. I, encara que en el joc prenguis partit de manera descarada, tu saps en el teu fur intern que tot és discutible, que tot és dubtós. Fins i tot saps que és molt probable que l'equip adversari tingui més raó que tu. Juges.

I ens direu: Què feu, tan ganàpies, que jugueu com criatures?

Us respondrem: Sí, som grans, però som vius! I, com que vivim, juguem! Ens autoorganitzem i apliquem a l'equip la nostra capacitat per al treball en grup. Exercitem els nostres dots de persuasió i de negociació. Espremem el nostre talent: tant el que la ciència i l'ofici ens han donat com el que ens han proporcionat els anys viscuts. Aportem al grup creativitat, i hi exercitem les nostres habilitats. Ens apliquem a aconseguir la derrota del contrari mitjançant refutacions erudites i elegants. Ens entrenem en la seducció del jurat a través de l'exposició clara i assertiva de les idees. Mentre interpretem un paper que no ens és propi, ens alliberem de complexos antics que no ens feien cap favor.

En el joc trobem una oportunitat d'atendre un repte intel·lectual. És un espai de comunicació i una ocasió de relacionar-nos. No només ens entretenim i ens divertim, sinó que, clarament, creem coneixement. Jugat és cultura.

Companys de Iubilo URV, estudiants sènior, escolteu aquesta crida: incorporeu-vos als equips de les edicions futures de la Lliga de Debat Sènior de la Xarxa Vives d'Universitats. Us ho recomanem. És segur que ho passareu molt bé.

A continuació, complementem aquestes conclusions amb uns relats personals de l'experiència de cadascú, en què no hem reprimit l'emotivitat. Un dels avantatges de l'ancianitat és que els vells podem parlar de sentiments sense cap empatx.

8. ELS RELATS DE CADASCÚ

La Lliga de Debat és un joc. En la mesura que proporciona plaer, el joc és una font de descans i relaxació. En qualsevol cas, el joc és una eina de transmissió de cultura i de valors, que ens prepara per a la vida i ens ajuda a socialitzar-nos. El filòsof i historiador neerlandès Johan Huizinga (1872-1945), autor d'*Homo ludens*, diu que el joc és una funció humana tan essencial com ho són el pensament o el treball (Huizinga, 1972).

En la fase de preparació del debat vam jugar de valent. Discutíem obertament, amb tota la força i la convicció, en una simulació estimulante i de creixement mutu. Manteníem discussions sobre conceptes clau (humanitat, intel·ligència, idees sobre el risc i l'amenaça...) i reflexionàvem sobre les posicions de diversos pensadors i científics a l'entorn del transhumanisme i el futur de la humanitat. Es pot dir que, més enllà del resultat final de l'activitat, el millor d'aquesta experiència ha estat el treball col·lectiu en aquestes sessions preparatòries. Podem recordar, per exemple, les lectures i comentaris dels textos que aportava el José Eugenio, les explicacions de l'Albert sobre el llibre *Sàpiens*, de Yuval Noah Harari (2014), o també les referències del món de la ciència-ficció i el cinema. Per fonamentar de manera sòlida els conceptes clau del debat, la Isabel va proposar acudir a la tradició aristotèlica i a d'altres pensadors clàssics perquè ens servissin de referents. La Carme estava molt pendent dels aspectes organitzatius i reglamentaris, i ens recordava com havien d'anar les sessions del debat i com calia ajustar bé totes les fitxes d'argumentació en un sentit o en l'altre. Quan la dinàmica de treball va impedir que la Maria i l'Estanis continuessin assistint a les sessions ordinàries, ells van seguir implicats en el projecte. Més

endavant, el Jordi es va afegir al grup i de seguida vam veure que la seva visió pràctica, clara i neta ens ajudaria a reconsiderar i complementar tota la feina feta fins llavors.

Com que els majors beneficiats de l'activitat lúdica són els mateixos jugadors, alguns dels impulsors i membres de l'equip de Iubilo URV que vam participar en la Lliga de Debat Sènior 2024 de la Xarxa Vives expressem seguidament, en primera persona, la nostra experiència.

Comencem primer amb el testimoni de Rosa Queral, que aporta la visió des de fora del grup dels participants, però que va ser un suport incondicional a l'assaig previ i a les sessions de la Lliga de Debat. Tot seguit, Jaume Aymí, capità de l'equip, aportarà el seu punt de vista, sobretot, dels dies de la competició. I, finalment, els competidors, la Isabel, la Carme, l'Albert i el Jordi, ens explicaran com ho van viure com a participants en tot plegat.

Gaudir del grup

«Es pot dir que, més enllà del resultat final de l'activitat, el millor d'aquesta experiència ha estat el treball col·lectiu en aquestes sessions preparatòries.».

RELAT PROPI DE ROSA QUERAL

L'ASSAIG PREVI

El dia 28 de maig, de 18 a 20 h, a la sala de juntes del campus Catalunya, els nostres companys i companyes de Iubilo URV que van formar l'equip per al Debat Sènior de la Xarxa Vives van fer un assaig per mostrar la recerca feta sobre el tema a debatre a la Lliga: «La intel·ligència artificial és una amenaça per a la humanitat?».

Els companys i les companyes havien treballat intensament, dedicant temps a la recerca d'informació, organitzant les dades i creant fitxes de lectura. Per aquesta raó, la Junta executiva de Iubilo URV va decidir organitzar un acte obert per compartir l'experiència amb la resta de la comunitat i el públic que ens segueix a través de les xarxes socials.

De cap manera podíem permetre que la informació i els coneixements adquirits quedessin només per al dia del debat; calia fer quelcom més. Per això, un assaig públic podria aportar beneficis com, per exemple, minimitzar possibles pànics escènics, intercanviar informació amb el públic al final de l'exposició, modificar comportaments o actituds, eliminar alguna informació poc rellevant i sobretot confirmar els coneixements adquirits gràcies a la consulta de nombroses fonts i materials treballats.

La convocatòria es va enviar per correu electrònic a totes les persones adherides i al grup d'Amics i Amigues de Iubilo URV, i el 25 de maig la vam publicar a la nostra pàgina web i a Facebook. Vam utilitzar per fer el cartell el mateix format de la Xarxa Vives.

L'assistència va ser considerable, i vam omplir la sala de juntes del campus Catalunya. L'equip, format per Jaume Aymí, Isabel Baixeras, Jordi Blay, Albert Bordons, Carme Crespo i José Eugenio García-Albea, va combinar l'exposició amb l'assaig. Durant el torn

de preguntes el públic va participar activament amb bones propostes que l'equip va recollir amb interès.

Com a comentari a una fotografia publicada a les xarxes, la Isabel va compartir una reflexió valuosa: «Els ancians de la comunitat URV analitzant, en grup, els reptes del futur». Un futur que potser no veurem, però desitgem, de tot cor, que les generacions que ens segueixen siguin capaces d'adaptar-se, d'organitzar-se, creant xarxes de suport, oferint oportunitats de descobriment i d'enfortiment mutu, compartint percepcions i experiències. Tot plegat ha de permetre que la creativitat no sigui anihilada, lluitant contra el cinisme com a estat crònic de desconfiança, antítesi de l'obertura necessària per enfrontar-se als reptes del futur i que entre totes i tots garanteixin i millorin la supervivència de la nostra espècie.

Per a Iubilo URV, el sentit de comunitat, les xarxes de suport i la creativitat són fonamentals. Aquest comentari inspira a considerar que, per a futures edicions, potser seria útil organitzar sessions separades: una d'informativa dirigida a tots els nostres públics i una altra d'assaig, amb persones expertes, per assolir objectius útils per al debat posterior, com ara el control efectiu del temps, el material, el format i els altres vessants de la comunicació oral i gestual.

Des de l'organització agraïm la gran resposta de totes les persones assistents i el compromís i esforç del nostre equip. Aquest esdeveniment demostra que el debat no només és una eina de reflexió, sinó també una oportunitat per construir comunitat, per créixer junts i afrontar reptes de futur amb esperança i determinació.

LA COMPETICIÓ

Els dies 5 i 6 de juny eren la data assenyalada per desenvolupar el debat. Els companys i companyes del nostre equip s'enfrontarien, per primera vegada, als equips de les altres universitats membres sènior de la Xarxa Vives. Havien treballat molt i estaven ben preparats, però disputar fil per randa una sèrie d'arguments sobre un tema tan nou i en un ambient no prou conegut era un repte afegit. Val a

dir que amb condicions diferencials de la resta de participants com a novençans, els nervis estaven a flor de pell.

Pocs dies abans de l'encontre, la Carme em va demanar que mirés d'agilitzar el tema logístic. Faltaven pocs dies i no sabien on anirien a dormir ni com seria el desplaçament cap a Barcelona. Amb Àngels Olivé vam aprofitar la visita al vicerector Jordi Diloli per esperonar-lo a buscar hotel pròxim al lloc del debat. Li vam ensenyar el que ja feia temps que s'havia publicat a la web de la Xarxa Vives. Després de mostrar-li el lloc de la UPC on es desenvoluparia el debat, va trucar a les companyes de l'Oficina de Compromís Social i es va posar la maquinària en marxa per trobar un hotel proper i que respongués als criteris marcats per la URV per a la reserva de places d'hotel. També es van fer les gestions per contractar un taxi gran per facilitar el transport.

El dia del debat vam marxar amb un taxi negre amb els vidres tintats cap a Barcelona. Vaig pensar que com a coordinadora de lubilo URV havia d'estar al costat de l'equip i intentar, en la mesura de les meves possibilitats, fer de cronista del que estava passant.

Ja havia estat al Campus Nord de la UPC. Durant un temps vaig estar de responsable de Salut Laboral de la Interuniversitària de CCOO i havia visitat algunes vegades la seu del Centre Nacional de Condicions de Treball (INSST), veïna del campus on s'havia de desenvolupar el debat.

Quan vam arribar a la sala, vam comentar amb la Isabel la ubicació gens bona que ens havien assignat. En entrar-hi tenien situats els ordinadors i les taules de so. El so no arribava al fons de la sala perquè no hi havia altaveus i se sentia molt poc del que passava a la part davantera, on havien ubicat l'escenari. Cables arraconats al costat d'un estret i únic passadís dificultaven el pas i una sèrie d'allargadors per terra que ens van permetre carregar les bateries dels mòbils. L'aula només tenia una porta d'entrada i sortida just al costat dels equips tècnics. Era evident que no s'havia pensat prou en l'espai. Es van haver de buscar altaveus perquè poguessin seguir el debat amb condicions les persones situades al fons de la sala.

El debat va transcórrer bé, tot i la repetició atzarosa del mateix posicionament de defensa. La presentació gràfica del nostre equip era austera, concreta, amb poca lletra i amb fotografies adients, com manen els cànons de les presentacions orals. Això va fer que la repetició del mateix argumentari no es fes feixuga, visualment parlant. Oralment, en cada repetició el nostre equip millorava el nivell de defensa introduint elements de debat nous.

Em va cridar l'atenció el cronògraf. Mentre marcava clarament el temps d'exposició, no feia el mateix amb el temps de resposta, que quedava a mercè del jurat sense que els concursants i el públic ho poguessin saber.

Durant tota la sessió dels debats, vaig anar fent fotografies i penjant-les a les xarxes socials. Vaig connectar amb el vicerector i li vaig enviar l'enllaç del debat en directe. També el vaig compartir amb la Junta executiva, els diferents grups de Iubilo URV i el vaig publicar a les xarxes.

Tot i no constar com a membre de l'equip, el tracte de la UPC amb mi va ser excel·lent.

Vam quedar empatats amb els segons i la defensa de la nostra Isabel va fer possible que ens emportéssim cap a casa un trofeu a la millor oradora. L'esforç fet per l'equip va ser compensat. Però el més bonic, al meu entendre, va ser el procés de la formació de l'equip: la preparació, les discussions, els neguits i, sobretot, la companyonia i l'amistat, vertaders pilars que hem de seguir potenciant per fer de Iubilo URV una comunitat on la gent importa.

El millor

«El més bonic va ser el procés de la formació de l'equip: la preparació, les discussions, els neguits i, sobretot, la companyonia i l'amistat».

De cara a noves edicions, s'haurien de tenir en compte les ubicacions, disposar de llocs segurs, amb bona acústica i que facilitessin la connexió dels ordinadors, la recàrrega dels mòbils, i un

cronògraf clar per a les respostes. A més, cal evitar que un mateix equip repeteixi sempre el mateix argumentari. Per la nostra part, caldria treballar amb més temps el tema logístic i de comunicació, i fer arribar a tot Iubilo URV les notícies sobre l'esdeveniment i l'enllaç que l'organització faciliti per seguir en directe el debat.

Assistir a aquest acte de la Lliga de Debat va ser tot un privilegi i crec que s'hauria d'aconseguir una audiència amb més públic.

RELAT PROPI DE JAUME AYMÍ

REPORTATGE SOBRE EL DEBAT

Després de tota la feinada dels darrers mesos, finalment va arribar el dia 5 de juny del 2024, data de començament del Debat. Es tractava de matinar de valent per arribar a temps a Barcelona, a les instal·lacions de la Universitat Politècnica de Catalunya, i amb bon ànim començar la nostra actuació.

Cal dir que l'aula on van tenir lloc els diversos actes de competició no era ni de bon tros prou adequada, especialment pel que fa a aspectes de seguretat, ja que només es disposava d'una única porta d'accés en unes instal·lacions d'un soterrani que, a més, no disposaven, almenys al principi, d'una acústica adient. És cert que es van resoldre ben aviat aquests problemes de so, i que després la UPC va explicar que en aquelles dates hi havia dificultats d'espai per causa d'altres esdeveniments oficials. Aquests detalls tècnics val la pena de consignar-los per si en alguna altra ocasió ens correspon a nosaltres l'organització d'aquest tipus de debat.

Les sis universitats participants van ser les següents:

- Universitat Autònoma de Barcelona
- Universitat de Barcelona
- Universitat Politècnica de Catalunya
- Universitat Pompeu Fabra
- Universitat Rovira i Virgili
- Universitat de Vic-Universitat Central de Catalunya

En funció d'aquesta participació (que en principi s'esperava que fos més nombrosa), havia tingut lloc prèviament un sorteig per organitzar la fase eliminatòria dels diversos enfrontaments entre els participants, de forma que totes les universitats dispo-

sarien de tres sessions diferents de debat entre si per arribar després per millor puntuació, si era el cas, al sorteig de la sessió final que clouria el debat. A nosaltres ens va correspondre, per aquest ordre, enfrontar-nos a la Universitat de Barcelona (matí del primer dia), a la Universitat Pompeu Fabra (final de la tarda) i a la Universitat de Vic-Universitat Central de Catalunya (segon dia al matí).

Ja ha quedat dit en algun moment d'aquesta publicació que la URV era principiant en aquest tipus d'activitat i, per tant, no coneixíem per dins l'ambient i l'estratègia. En canvi, observàvem de seguida un bon nivell de coneixença i de companyonia entre els altres diversos equips participants. A més, no sabíem com jugaria o no al nostre favor el fet d'haver estat acceptats com a equip participant no sent alumnes oficials sènior, sinó personal jubilat de la mateixa URV. És cert que després vam confirmar clarament que molts dels nostres contrincants eren també personal jubilat i amb una llarga trajectòria professional als seus centres. I d'altra banda, vam veure que els equips disposaven d'assessorament tècnic que els havien preparat adequadament (vam poder conèixer personalment alguns d'aquests assessors). El coordinador del Debat de la Xarxa Vives comunicava el resultat als capitans de cada equip, i explicava quins havien estat els punts forts i febles de cada equip durant el debat. Els criteris que consten a les actes oficials rebudes són: presentació d'una tesi sòlida, qualitat del discurs, flexibilitat i capacitat del debat, demostració del domini del tema, aspectes formals de l'exposició i posada en escena i, finalment, l'actitud. En la nostra opinió, cal dir que el sistema de puntuació i la seva aplicació pràctica no eren prou entenedors, i que caldria clarificar-los i simplificar-los en edicions futures.

Un altre punt important a destacar és que per sorteig es coneix, cinc minuts abans de la sessió, quina posició correspon a cada equip, a favor o en contra de l'enunciat. Lògicament, pot correspondre en cada cas la mateixa opció o cada vegada canviada: en aquest cas, les nostres tres actuacions van ser per defensar sempre el sí, és

a dir, que efectivament la intel·ligència artificial és una amenaça per a la humanitat. Pensem també que aquest punt s'hauria de corregir en el futur per poder observar una millor variació i complexitat en el treball dels equips, i així enriquir la dinàmica dels diversos enfrontaments.

A millorar

«Les nostres tres actuacions van ser per defensar sempre el sí, per sorteig. Pensem que s'hauria de corregir aquest punt per poder observar una millor variació i complexitat en el treball dels equips».

Més enllà de les opinions i vivències personals dels companys del nostre equip, correspon aquí ressenyar amb un cert detall les puntuacions que vam obtenir en les tres sessions.

1. Sobre un valor màxim de 5, vam obtenir un 3,66 total enfront d'un 3,78 de la Universitat de Barcelona. La diferència que recull l'acta és a causa del 0,25 de penalització que vam patir per haver-nos excedit uns segons en l'administració del temps. És a dir, vam pagar la nostra inexperiència a l'hora de pautar estrictament el temps de les intervencions, de manera que sense aquesta penalització hauríem guanyat aquest primer debat. No es pot dir que no estiguéssim avisats, però ens calia potser més pràctica en aquest punt.
2. En la segona intervenció, davant la Universitat Pompeu Fabra, ens va passar una mica el mateix, ser penalitzats amb 0,25 pel límit del temps, però aquesta vegada la nostra millor puntuació global de 3,92 enfront de 3,89 ens va permetre encarar el segon dia amb optimisme.
3. A la darrera intervenció, aquesta vegada contra la Universitat de Vic, ja vam aprendre la lliçó i vam poder quadrar bé els temps en tot moment. Van ser els rivals aquesta vegada els penalitzats pel mateix motiu, i la puntuació final al nostre favor va ser de 4,01 contra 3,78.

Davant aquesta situació ens quedava el dubte sobre si accediríem o no a la final segons la puntuació d'altres universitats. Ens va semblar per uns moments que sí, de manera que estàvem pendents de poder confirmar aquest fet al nostre vicerector perquè vingués a presenciar la final i el possible triomf. No va ser així, a migdia vam saber que quedàvem tercers.

En tot cas, com a confirmació final de tota la nostra actuació com a equip, vam ser premiats amb el guardó a la millor oradora de tot el debat per a la nostra companya Isabel Baixeras. Podeu veure algunes imatges dels dies de la competició al capítol 14.

Dies després vam poder respondre una enquesta que la Xarxa Vives va enviar als equips participants. Vam comentar, en general, molts aspectes positius que ens havia aportat aquesta experiència, especialment perquè pensem que és un bon al·licient per mantenir-se actiu i treballar en equip, i també perquè millora la nostra capacitat intel·lectual i el benestar psicològic general. D'altra banda, vam poder fer constar alguns aspectes de millora, especialment referits a la necessitat de revisar la mecànica del debat: ampliar el nombre d'intervencions, modificar les penalitzacions per temps, revisar el sistema de preguntes i refutacions. Vam tenir finalment l'ocasió de suggerir diversos temes possibles de debat per a l'edició 2025.

RELAT PROPI D'ISABEL BAIXERAS

FASE DE PREPARACIÓ

Un matí de febrer del 2024 travessava l'aparcament del campus Catalunya meravellada de tant d'ordre. El paviment, que llis! Els cotxes, que ben posats: tots arreglats! Fa uns vint anys, els universitaris deixaven els vehicles de qualsevol manera al pati de l'antiga caserna i s'hi generaven sovint unes bregues (que si m'has taponat, que si no em deixes passar...) que alguna vegada havien arribat a l'oficina de la Sindicatura de Greuges de la URV, oficina que va ser al meu càrrec en el període —que avui queda lluny— comprès entre el 1999 i el 2004.

Dispenseu. Les jubilades sempre expliquem batalletes que no interessin a ningú. Vaig al gra. El cas és que, impressionada per la pau de l'estacionament, vaig pujar al pis quart per l'escala de la porta quarta i vaig entrar per primera vegada a la seu de Iubilo URV, a una reunió que havia convocat el Jaume: una persona amb qui, fins aquell moment, només havia parlat per telèfon i de qui coneixia, només, el seu bonic accent ebrenc. Hi havia més gent. Vaig saludar antics amics —la Carme, l'Estanislau— i vaig conèixer la resta de membres de l'equip amb qui hauria de concursar a la Lliga de Debat Sènior de la Xarxa Lluís Vives: el José Eugenio, l'Albert, la Maria...

Jo no les tenia totes. Veia que no estaria a l'altura dels meus companys, tots persones erudites. Altrament, el tema de debat de la Lliga 2024 m'inquietava: «La intel·ligència artificial és una amenaça per a la humanitat?». Jo no en sabia res, de la IA! Si havia acceptat la invitació del Jaume és perquè m'agrada posicionar-me esportivament en relació amb qualsevol idea i jugar a defensar-la amb totes les forces mentre, per dins, dubto de tot el que dic. Però era clar que amb aquesta actitud entremaliada no n'hi hauria prou, i que el compromís de representar la URV en una competició de la Xarxa Vives m'obligava a ser seriosa i responsable.

En qualsevol cas, tant a la primera reunió com a les successives l'afluència constant d'idees al voltant del tema em va fer veure que l'aventura oferia molt bones perspectives. De seguida em van agradar tots els membres de l'equip —tots i totes!—, vaig estar contenta de formar-ne part i, amb l'ajuda de tothom, em vaig començar a endinsar en els conceptes d'*intel·ligència artificial*, d'*amença* i, fins i tot, d'*humanitat*.

Poques setmanes més tard l'equip es va reestructurar lleugerament. L'Estanis i la Maria no podien continuar per raó dels compromisos respectius (els jubilats estem, sovint, molt enfeïnats!) i els va substituir el Jordi, que amb la mirada neta i crítica que té qui s'incorpora de fresc en una discussió, va posar ordre a l'olla de grills i va centrar les coses.

Per ser més eficaços, ens vam repartir els dos posicionaments possibles: resposta afirmativa i resposta negativa. Jo no tinc res contra la IA, però per comoditat em vaig oferir per al discurs alarmista. Sempre és més fàcil ser retrògrada i apocalíptica que no pas ser confiada i valenta. Ara: com que la pregunta era capciosa, ja es veia que, en el joc que proposava la Xarxa Vives, allò que més comptaria seria la brillantor dels arguments, l'estètica de les il·lustracions, l'assertivitat en l'exposició i l'originalitat en el discurs.

ELS DEBATS

La Carme, el Jose Eugenio i jo ens havíem preparat per sostenir que la intel·ligència artificial és, efectivament, una amenaça per a la humanitat. El Jordi i l'Albert s'havien preparat per negar-ho. En el sorteig de la primera sessió ens va tocar defensar que sí. Els tres «especialistes en el *sí*» vam anar sortint a l'estrada amb les nostres flamants fitxes de cartolina i, si bé ens vam passar una mica del temps assignat, no vam fer mal paper.

Tot feia suposar que, en el sorteig de la sessió següent, l'atzar es decantaria i ens tocaria el posicionament pel *no*, i que l'Albert i el Jordi desplegarien l'argumentari que duien tan ben preparat.

Decepció! Va tornar a sortir el *sí*! Aquesta vegada, els ponents vam omplir de gargots les nostres precioses fitxes, vam improvisar unes argumentacions renovades i vam sortir a escena bastant indignats. Es van acabar els somriures al jurat i les consideracions amables amb l'equip contrari. Se'ns va posar un timbre metàl·lic a la veu. Vam respondre impacientment les preguntes dels adversaris. Com que l'Albert i el Jordi, quiets als seus seients, tenien la mosca al nas, van interrompre amb mala bava els discursos dels contraris. Primmirats com en un tribunal de tesi, no en deixaven passar ni una. Amb el mal humor que exhibíem tots plegats, el prestigi del nostre equip va pujar com l'escuma.

Al final de la tarda, en comptes d'anar a descansar per estar l'endemà ben reposats, ens vam reunir en comitè de crisi amb el Jaume, el nostre capità, i ens vam plantejar seriosament que res no impediria l'atzar d'adjudicar-nos, de nou, el posicionament pel *sí*. Vam considerar unànimement que aquesta eventualitat hauria privat els millors valors del nostre equip de participar en el debat. Calia adoptar solucions d'emergència. La decisió va ser que, aquella nit mateix, el Jordi i l'Albert prepararien tot l'argumentari del *sí*, i l'endemà, fos quin fos el posicionament que se'ns adjudiqués, serien ells qui representarien la URV en el debat.

El resultat del sorteig del dia següent va ser el mateix de sempre: la URV havia de defensar el *sí*. Ja no ens vam indignar... Vam moure els llavis lleument, amb el somris resignat de qui ja sap que la ruleta no l'estima però segueix jugant-s'ho tot. El cas és que els «especialistes en el *no*» van actuar en favor del *sí*, i van desplegar els seus dots de convenciment amb l'assertivitat accentuada de qui defensa amb passió abrandada allò que no creu. Amb allò que se'n diu la fe del convers. Per tant, vam tornar a quedar bé. Molt bé.

Bona tasca, bon paper

«Per tant, vam tornar a quedar bé. Molt bé».

ELS PREMIS

Tot i que érem «clarament» els millors, el cert és que no vam passar a la final. Empatats amb l'equip que havia quedat segon, vam ser considerats tercers.

Tot i això, vam tenir un premi de consolació. El jurat va decidir per unanimitat nomenar-me la millor oradora. Que contenta que em vaig posar! Me'n vaig alegrar per molts motius, i crec que en puc explicar alguns:

- a. Perquè estàvem estressats i, amb aquella tensió, els sentiments es desboquen una mica.
- b. Perquè el meu complex de noia tímida i poqueta cosa m'havia fet sempre descartar un èxit d'aquest tipus.
- c. Perquè el meu equip havia fet un treball magnífic, exquisit, resultat de la pluridisciplinarietat dels seus membres i de moltes hores de treball intens. El nostre discurs era rígid i creatiu. Dúiem totes les argumentacions científicament justificades amb les fonts corresponents. Érem tan bons, almenys, com els guanyadors. El premi al millor orador va ser, en realitat, un reconeixement al treball col·lectiu. Vaig actuar en més ocasions, certament, que els altres, però això va ser perquè la fortuna ho va voler així. La meua presència repetida davant del jurat em va donar més ocasions de lluir. És més: la indignació que duia a sobre pel sentit reiteratiu dels sortejos em va dotar d'una impertinència i una arrogància verbal que normalment reprimeixo.
- d. Perquè entre el 1999 i el 2004 vaig exercir les funcions de síndica de greuges per a la Universitat Rovira i Virgili. Va ser una feina compromesa, delicada i gens fàcil, que mobilitava a mantenir distàncies prudents tant amb els queixosos com amb els implicats en les queixes. El càrrec té una durada breu, perquè no és convenient que s'allargui en el temps. És per això que em fa contenta que els meus interlocutors

de fa 25 anys —molts d'ells companys i companyes de Iubilo URV— comptin amb mi com una jubilada més de la universitat que estimo, i que em mostrin una amistat que endevino sincera.

- e. Perquè, així que vaig tenir a les mans el diploma, la Rosa, satisfeta, en va enviar una foto al vicerector de Compromís Social i Sostenibilitat, i perquè em va omplir d'orgull tornar a Tarragona amb un diploma de la Xarxa Lluís Vives per a la meva universitat.
- f. Perquè Iubilo URV ha penjat el meu diploma a la paret del seu despatx: el de la planta quarta de l'escala quarta del campus Catalunya.

RELAT PROPI DE CARMÉ CRESPO

RECORDS AGREDOLÇOS D'UNA LLIGA DE DEBAT

Recordo el primer dia que el Jaume ens diu a la Junta de Iubilo URV que el vicerector li ha demanat que siguem nosaltres qui busquem estudiants sènior a la nostra universitat per anar al Debat Sènior de la Xarxa Vives 2024.

La meua primera reacció va ser una mica visceral (com soc jo...), d'enuig, i els comento que la URV no té pràcticament estudiants sènior, ja que no té cursos reglats ni programes d'ensenyament sènior com tenen altres universitats, i ja no dic les grans, sinó del nostre nivell: Lleida, Girona, etc. Comento que jo mateixa vaig explicar, a diversos vicerectors anteriors i a l'única rectora que vam tenir, com jo mateixa estudiava Història de l'Art a la Universitat Sènior de la UB amb gent que ni tan sols tenia cap carrera o era d'altres especialitats, com jo, i era un èxit total de convocatòria. Tot i això, dic que em sembla bé i cal buscar aquests estudiants sènior perduts entremig de tots els graus de la URV, la qual cosa va esdevenir *missió impossible*, parafrasejant la pel·lícula.

En vista del nul èxit, el vicerector demana a la Xarxa si poden ser els membres de Iubilo URV (que som sènior, també, però no estudiants) qui hi participem, i ens diuen que sí.

Un cop arribat en aquest punt, ara toca veure qui hi vol participar, i aquí tenim una altra travessia pel desert. La gent no pot o li afronta el tema proposat: «La IA és una amenaça per a la humanitat?». Aquí m'hi apunto per solidaritat, però amb la intenció que si hi ha molta gent, em retiro; no va ser així, i aquí em veus a mi, que no m'agrada parlar en públic, i que a més tenia un problema familiar, no gaire favorable a poder-hi participar, pujada al carro del debat.

He de dir que hi van haver moments de tot, alts i baixos. No tot van ser flors i violes. Les desavinences amb el Jaume per temes d'interpretació de reglament i modus operandi del debat van ser antològiques; també els últims dies vam tenir problemes amb la logística, fins que vaig demanar ajuda a la Rosa (coordinadora de l'ubilo URV), qui també va fer la comunicació al minut dels debats per a tota la comunitat URV. La sang, però, no arribà al riu i vam aconseguir al final un grup viu i compacte tot i venir de mons diferents de la URV.

Dificultats inicials

«Hi va haver moments de tot, alts i baixos. No tot van ser flors i violes. Al final vam aconseguir un grup viu i compacte tot i venir de mons diferents de la URV».

A l'hora de repartir-nos qui preparava la defensa del *sí* i del *no* (documentació, elaboració de fitxes, etc.) em toca defensar que la IA *sí* que és una amenaça per a la humanitat, quan jo en aquell moment estava lluny de creure això per motius varis, perquè a la meva feina al Servei de Recursos Científics havia vist com ja s'utilitzava quelcom semblant en aplicacions en biomedicina, etc. i també el meu fill, doctor en astrofísica que estava fent recerca vers la Via Làctia, havia utilitzat el big data a la tesi i això els va permetre fer el mapa estel·lar actual de la Via Làctia en 18 mesos amb una troballa (*baby-boom* estel·lar de fa 4.000 milions d'anys) (Mor et al., 2019) inclosa. Tot i això, suposeu que el meu fill em va ajudar en la recerca? Doncs no! Quan ja vaig tenir recollits els arguments del *sí* li vaig ensenyar les fitxes i em va fer tot un assortiment d'arguments en contra que ens van servir als del *no* per refutar el *sí*. Què hi farem! Sempre estimarem els fills.

Pas següent: ja tenim les fitxes, ara cal muntar els argumentaris per al debat, i en principi preparem l'Albert el *no* i jo mateixa el *sí*. El de l'Albert, suau i bo com correspon al *no* és una amenaça; el meu, els companys el classifiquen de catastrofista! —cosa normal, perquè

tot era dolent, començant per la bomba atòmica i Oppenheimer. Per cert, després d'aquesta recerca la meua posició respecte a la IA va canviar lleugerament.

Un cop aquí, demanem a la Isabel que en faci una altra a partir de les fitxes i així entre totes dues ho encarem; carai, però és que el cap de la Isabel és un terratrèmol, amb milers d'idees que li sorgeixen cada dia i que s'han de concentrar en pocs minuts i ella escriu que escriu (crec que va escriure sis propostes diferents o més, que ja n'he perdut el compte), i així em vaig trobar 24 hores abans del dia D agafant tots els escrits de la Isabel, documents i idees i concentrant-los en els quatre minuts d'exposició inicial que em tocava defensar a mi. La resta li van servir i molt a ella per a les brillants refutacions posteriors que va defensar i que li van fer guanyar el premi a la millor oradora dels debats, ja que no tenia rival possible. A YouTube em remeto si no em creieu. Hi podeu veure les gravacions dels debats, que hem referenciat al darrer capítol, el 14. Més sort van tenir les conclusions del José Eugenio, que van ser tan encertades que no calia canviar-les gaire.

De les competicions ens va tocar defensar el *sí* en tots els debats, la qual cosa no ens va agradar gens, ja que el *no* de l'Albert i el Jordi el tenien súper ben preparat. Després de la mala sort del primer dia (repetició del *sí* matí i tarda), i sort que la Isabel utilitzava arguments diferents en els torns de refutacions. Per cert, no he dit que en el primer debat vam perdre pels segons de més afegits i el segon debat el vam guanyar perquè els nostres arguments eren millors i ja no ens van agafar de *pardillos* amb el temps, ja que sembla que el temps penalitza més que els mals arguments.

De camí tot caminant a l'hotel vam pensar que si el dia següent ens tornava a tocar el *sí* no podríem suportar-ho, i que havíem de canviar de ponents i d'arguments; així, després de sopar, en lloc d'anar a descansar, doncs apa, els passem el relleu a l'Albert i al Jordi, que refan el relat a les quartilles en blanc o a les mateixes pel darrere; per si de cas, la mà del Jaume tampoc encerta aquest cop el *no*, tal com va ser. L'Albert i el Jordi ho fan genial i, això sí, ben

acompanyats de la nostra refutadora perpètua en tots els debats, que per això sabíem que ja s'emportava el premi (és broma).

I un cop acabat el debat, contents, marxem cap a Tarragona i a la Isabel li agafa un gran sentiment de culpabilitat perquè diu que no competim lleialment, perquè nosaltres som ex URV i que els altres són alumnes. No puc veure-la així, i de retorn cap a casa busco els components de les altres universitats; heus aquí que són professors d'universitats, d'instituts i algun exdirector general de Medi Ambient. Aquí tanquem el tema culpabilitat.

La meva conclusió de tot plegat és que ha estat una experiència única, diferent de tot el que jo havia fet fins al moment, que si bé vaig tenir dubtes a l'inici, m'hi vaig abocar de ple. També crec que puc dir que el grup de Iubilo URV ara som un grup d'amics, una amistat de sèniors, però no acabats, generada per la feina intensa en poc temps, però engrescadora, i per aquesta raó deixo el meu lloc perquè al proper debat altres dones i homes del nostre col·lectiu provin la pujada d'adrenalina que s'experimenta amb l'experiència.

RELAT PROPI D'ALBERT BORDONS

UNA MOLT BONA EXPERIÈNCIA

La meva sensació resumidament és que m'ho he passat molt bé, he après força i he fet bons amics.

M'HO HE PASSAT MOLT BÉ

Com a professor emèrit, reconec que faig el que vull i que intento passar-m'ho bé. Segueixo fent algunes «feines» amb alguna utilitat per a la URV o altres institucions, i com a difusió de la ciència en general, i sobretot pel plaer de conèixer i per curiositat científica, en el meu cas sobretot biològica, però també geogràfica, històrica, antropològica i lingüística. L'experiència dels anys fa que pugui fer algunes classes, tutories, avaluacions, xerrades, informes, revisions, articles divulgatius i altres «feines» relativament bé i gaudint-ne.

M'ho he passat molt bé en la preparació del debat i la competició, perquè ha estat una «feina» més d'aquestes que m'agraden, però sobretot perquè ha estat una feina de grup, d'un equip on m'he sentit integrat des del primer dia, i on he gaudit de moltes hores de converses, discussions, coincidència o no d'opinions, però sobretot de riures, ocurrencies i bones emocions.

HE APRÈS FORÇA

He après sobretot del tema de debat, o sigui, de la intel·ligència artificial —la IA—, encara que a hores d'ara segueixo sense ser-ne un expert ni un usuari habitual. Donat que el tema de debat era «La intel·ligència artificial és una amenaça per a la humanitat?», vaig

començar a estudiar el tema, i a banda de *Wikipedia* i Viquipèdia, vaig llegir força articles i alguns llibres, dels quals recomano sobre tot el de [Melanie Mitchell \(2024\)](#). Durant les setmanes de preparació, tot l'equip vam anar submergint-nos en el coneixement de la IA, en els riscos, per un costat, i els usos i beneficis, per un altre. Tot plegat, en arribar a la competició amb els debats dominàvem força el tema, i la prova va ser que em sembla que vam mostrar més coneixement i moltes més fonts d'informació —les evidències, en deien— que molts dels altres equips.

Bona preparació

«En arribar a la competició dominàvem força el tema, i em sembla que vam mostrar més coneixement i moltes més evidències que els altres equips».

Aquests mesos també he après a debatre, i sobretot a preparar un debat. La prova evident és que malgrat estar convençut del «no-amença» i tenir-ho tot preparat per a aquesta posició, al darrer debat, on ens va tocar per tercera vegada haver de defensar el «sí-amença», vaig actuar en aquest sentit, preparant-ho la nit abans, amb el Jordi també canviant de posició. Diuen que ho vam fer prou bé, amb bons arguments, i amb la brillant companya Isabel —destacada refutadora—, que merescudament va guanyar el premi de millor orador/a. Per saber debatre, encara que estic d'acord en les tècniques de comunicació que vam anar aprenent, per mi el més important segueix sent el coneixement del tema.

HE FET BONS AMICS

Encara que els coneixia quasi tots, la relació prèvia era molt escadussera. Però aquests mesos les interrelacions han estat contínues i m'ha sorgit un veritable afecte amb tots. Hem passat moments molt entretinguts i força estones divertides. Els sis membres de l'equip que hi hem estat implicats més a fons (la Isabel, la Carme, el Jau-

me, el Jordi, el José Eugenio i jo) hem conviscut moltes hores. En anar-nos coneixent, crec que la confiança i la sinceritat han anat augmentant, i a hores d'ara els considero molt bons amics.

Tots i cadascun dels sis tenim les nostres particularitats i uns caràcters i comportaments força diferents. Lògicament, això ha fet que les relacions entre alguns no hagin estat uniformes, i hi ha hagut alguns breus moments de discussió o malestar. En el meu cas, sobretot al principi, m'irritava que el Jaume no contestés els correus electrònics, tal com un dia li vaig manifestar, un pèl empenyat. Crec que fins i tot aquestes petites desavinences han contribuït a reforçar l'amistat.

Per mostrar com som de diferents tots i per la confiança que tinc amb els altres cinc, m'atreveixo a fer...

UN PETIT JOC

Suposem que en aquest relat haguéssim decidit no incloure els nostres noms. D'aquesta manera potser hauríem utilitzat pseudònims per citar-nos cadascun. Doncs si em toqués a mi designar el de tots, se m'ocorren els següents, a partir d'alguns trets del caràcter o comportament durant aquest temps de convivència:

n.	Pseudònim	Acrònim
1	Esforç Neguitós	ESFONE
2	Manuscrit Revisor	MANURE
3	Gaudi Brillant	GAUBRI
4	Ordenació Digital	ORDEDI
5	Coneixement Garantista	CONEGA
6	Pragmatisme Suggestor	PRAGSU

L'ordre és a l'atzar, i, com veieu, són de dues paraules, com el nom i cognom, on tant un com l'altre —independents l'un de l'altre— a mi em recorden la persona. En aquests noms el gènere gramatical és independent del gènere sexual. Segur que l'opinió

de cadascú d'ells cinc sobre aquests pseudònims pot ser diferent, i segur que descobriran fàcilment qui som cadascú, però convido els altres amics i coneguts que ens identifiquin, a veure si saben trobar-nos a tots.

ALGUNA ANÈCDOTA

Jo no havia utilitzat mai abans l'aplicació d'IA ChatGPT, però just els dies de la competició me la vaig instal·lar al telèfon i vaig començar a fer-li preguntes. El darrer dia després de la competició, on vam quedar tercers com a equip, li vaig preguntar quin equip havia guanyat la Lliga de Debat Sènior 2024, i em va respondre que la UPC, cosa no certa. Tot seguit, li vaig dir que s'equivocava i em tornà a respondre: «La Lliga ha estat guanyada per la UPF», tampoc l'encertà. Li vaig tornar a dir que s'equivocava, i em digué: «La Universitat Rovira i Virgili ha guanyat la Lliga de Debat Sènior de la Xarxa Vives 2024»: Fantàstica resposta!

El dia següent per fi la va encertar: va guanyar la UB, llàstima!

RELAT PROPI DE JORDI BLAY

«JO, DEBAMENT SOBRE INTEL·LIGÈNCIA ARTIFICIAL?
TU HAS FUMAT ALGUNA COSA!»

Aquesta va ser la primera resposta a la proposta de participar en el grup de debat sènior de la URV, si no recordo malament feta per la Carme un dia que anàvem a una de les sessions de classe d'italià organitzades per Iubilo URV. Com a jubilat recent (setembre del 2023) em semblava fantàstic el curs d'italià, una activitat que de feia temps volia fer. En realitat, en aquell moment m'havia proposat de fer només aquelles coses que em vinguessin realment de gust, que per a això ens jubilem, no?

No em semblava que debatre sobre intel·ligència artificial fos una cosa gaire atractiva. Primer, perquè no tenia gaire idea sobre el tema i pensava que poca cosa hi podria aportar. Segon, perquè no soc gaire de debatre. Si hi ha cap discussió que s'allarga mínimament, opto per donar la raó a l'altre ràpidament amb algun *potser tens raó*, tot i que internament segurament continuo pensant el mateix que abans, i potser més. Aquest segon argument va ser el que vaig fer veure a la Carme: «Si em poso a debatre jo, segur que perdem», li vaig dir. «A la mínima donaria la raó als contrincants, per no discutir...» La Carme em va deixar per inútil, suposo. No vaig saber gaire cosa més del tema fins que passat Nadal el Jaume em va trobar pel carrer i em va insistir. Li vaig dir el mateix: «No en sé, del tema. Resposta: «Jo tampoc, home; d'aquí al dia que es faci el debat segur que en saps prou». Vaig treure el segon argument: «És que no sé debatre. Com a molt puc preparar arguments». Resposta: «D'acord, fitxat! Dimarts que ve a la Facultat a les sis de la tarda. Adeu». Ostres!

El dimarts següent vaig arribar al despatx de Iubilo URV a l'hora indicada. Coses del destí, era just el despatx que havia ocupat jo

mateix anys enrere com a professor quan els de Geografia estàvem ubicats al campus Catalunya. D'entrada, ja em trobava com a casa... I allà em vaig trobar un quintet curios: la Carme, la Isabel, l'Albert, el Jaume i el José Eugenio. I, després de saludar-me, es posen a debatre sobre si tenien prou arguments a favor o en contra. «No sé si podré aportar gaire cosa», vaig pensar, «em porten tres mesos d'avantatge». Vaig estar a punt de dir: «Bé, ha estat un plaer, ja m'ho explicareu...». Em devien veure les intencions, perquè aleshores l'Albert em va comentar: «Pots repassar les fitxes que hem fet amb arguments a favor i en contra, a veure què et sembla». Desenes de fitxes fetes a partir d'articles, llibres, notícies, vídeos i altre material sobre les virtuts i els defectes de la IA. Mare de Déu. Per acabar-ho d'adobar, el Jaume em va dir: «Mira't aquests documentals i aquestes pel·lícules de ciència-ficció que aborden el tema, segur que en pots treure idees». Malament, no hi havia *Alien*. Vaig pensar: «Em quedaré una estona perquè no es digui, i ja veurem si torno...».

Però la capacitat de persuasió del quintet de sèniors devia ser molt alta, perquè no només vaig continuar anant a les reunions, sinó que vaig repassar les fitxes una a una i vaig mirar els vídeos, incloent-hi també alguns de la competició de la Lliga de Debat per a secundària, per veure com anava la cosa. I al cap de 15 dies jo ja estava prou avançat: els companys havien fet una tasca ingent! Les fitxes sintetitzaven perfectament cadascun dels fets que permetien construir arguments a favor o en contra de la IA: hi havia resum i explicació de cada fet o opinió, referència de l'autor i l'obra, i possibles arguments per a rebatre-ho. Impecable. I els vídeos i pel·lícules em van fer veure que això de la IA era realment interessant com a camp d'aplicació i possibilitats futures, i com a element susceptible de filosofar per pensar sobre el futur de la humanitat.

En diverses reunions vam comentar què en pensàvem, de tot allò. I el cert és que en general no crèiem que s'hagués de pensar en la IA més com una amenaça que com una oportunitat, tot i els perills potencials. «Segur que es podrà controlar i donarà molts beneficis», pensàvem. «Però hem de determinar arguments a favor i

en contra. Si ens toca, com defensarem l'argumentació del *sí*, si no ens ho creiem cap de nosaltres? Ja ho veurem, ara fem l'exercici que cal fer per al debat!»

Vaig donar un cop de mà en la construcció dels arguments (no gaires, ja estaven bastant construïts) i especialment en la forma d'exposar-los en el debat. Recordo especialment les sessions en què les diferents versions de la introducció, les refutacions o les conclusions s'exposaven per part dels companys: «Que si llegim, que si no llegim. Com hem de llegir si es tracta d'un debat? Cal ser més natural. Cal que els arguments a favor o en contra siguin més clars: enumerem-los un, dos, tres... No, que no queda bé. Controla el temps, que et passes! Però com voleu que digui tot això en quatre minuts? L'expressió corporal, cal cuidar-la! No anem amb papers, anem amb cartolines. Per què? D'acord, però que tinguin el nostre logo darrere. I aguanta-les amb la mà esquerra, ho van dir al curset. Hem de fer un PowerPoint per a cada intervenció! Però com hem de fer un PowerPoint per a una refutació si no sabem què hem de refutar?» Etc.

Jo vaig intervenir sobretot en la introducció i la conclusió dels arguments del *no*. Però la veritat és que l'Albert ja ho havia treballat força. Crec que es va fer un bon exercici d'estructuració i síntesi dels arguments, tant en un posicionament com en l'altre. Teníem bones introduccions, conclusions i refutacions dels arguments. I refutacions de les refutacions. I en algun cas refutacions de les refutacions de les refutacions. I l'exposició estava ben plantejada, fent-ho cadascú amb el seu estil. I ajustada al temps! La Carme, la Isabel i el José Eugenio intervindrien per al *sí*, l'Albert i el Jordi, per al *no*. Res no podia sortir malament.

Molt preparat

«Teníem bones introduccions, conclusions i refutacions dels arguments. I refutacions de les refutacions. I en algun cas refutacions de les refutacions de les refutacions».

En un altre capítol ja s'explica com va anar la competició. En el meu cas, em van semblar molt interessants les diferents formes de preparar la Lliga per part dels diversos equips. Els de la UPC havien llegit —o ho semblava— diversos informes i dossiers sobre la IA i ho deixaven anar cada cop (pobres, haver de llegir tot allò). Els de la Pompeu havien muntat un xou audiovisual amb cartolines amb fotos que treien en el moment que l'orador exposava segons quin fet (caram, vols dir que quedarem avorrits amb els nostres senzills PowerPoints?) I els de la UVic vinga a dir que eren de Manresa i fent anar la sèquia amunt i avall (i si diem alguna cosa de la Tarragona romana?). N'hi havia que duien *coach* (i això els ho paga la universitat?). Ho havíem d'haver preparat diferent? (No, que en sabem més que ningú, caram!). Per mi, els moments memorables són els següents:

- 1) Quan ens diuen que perdem per excedir-nos de temps tot i fer-ho de conya al primer debat. Serem *novatos*!
- 2) Quan ens diuen que tornem a perdre per passar-nos de temps al segon debat, tot i que ho havíem brodat, amb la Isabel cada cop més convençuda de la maldat intrínseca de la IA i dels que la controlen. Agh! I sobretot quan ens diuen que no, que s'han equivocat els del jurat i ho han comptat malament! No m'estranya, amb el sistema de puntuació que fan servir...
- 3) Quan a la nit, per comptes de descansar, ens vam posar a escriure un discurs diferent del que ja teníem perquè, si ens tocava defensar el *sí* altre cop, repetir el mateix discurs segur que penalitzava, per avorrits. Qui m'havia de dir quan em vaig jubilar que tornaria a treballar de nit per tenir fet el treball a l'endemà...
- 4) Quan el dia següent cridem un taxi i, com és habitual, el taxista no coneixia l'adreça. El taxista: «No es preocupin, ara li pregunto al mòbil pel campus de la UPC i segur que arribem bé». Vam desconfiar tots de la IA quan el taxista la va fer ser-

vir per trobar la nostra destinació a la UPC i ens volia deixar a l'altra punta del campus, dient-nos que hi havia arribat. Un altre argument a favor del *sí*, vam pensar...

- 5) Quan al tercer debat ens torna a tocar defensar el *sí*, vam deixar anar tots els *mecàgums* possibles. No els poso aquí perquè no s'avindria amb l'estil de la publicació. Ho vam fer de conya i vam passar per damunt de la sèquia de Manresa a tota velocitat. Esperances d'arribar a la final...

Doncs no. Oh... «Potser s'enfadarà el rector. I el Diloli segur que no ve al lliurament de premis.» Però hi ha una bona notícia: la Isabel guanya el premi a la millor oradora. Tots molt contents! I és que, com ens van comentar alguns col·legues d'altres universitats, la Isabel en debatre es menjava amb patates el contrincant. I això que ella no veia gens clar això que la IA fos una amenaça. Si s'ho arriba a creure...

La competició va acabar i els equips van tornar a casa després del lliurament de premis. Què ha quedat de tot això? Per a mi, diverses coses. La primera, els bons moments passats amb els companys de l'equip de la URV. El cert és que ens ho hem passat d'allò més bé. Amb nervis i amb esforç, però també amb estones més relaxades i rient del que es podia riure. Crec que podem dir que n'han sortit unes bones amistats.

La segona cosa, veure que això de classificar-nos com a sèniors o tercera edat de vegades no sé si té prou sentit. La sensació és que, preparats com anàvem, ni que ens haguessin posat davant equips de professors universitaris, hauríem fet igualment un bon paper. Deu ser això, la vellesa activa, tu.

La tercera cosa —qui m'ho havia de dir—, constatar que això de la IA no és tan llaunot com semblava abans de començar. Hi ha aspectes de la IA que m'han semblat molt positius, i d'altres, certament preocupants, però en tot cas crec que és un tema interessantíssim i cada dia miro les notícies a veure quina cosa en diuen. Positiva o, més sovint, negativa. Tanta informació i no tinc clar com evolucionarà tot això. Potser que ho pregunti al ChatGPT...

9. IUBILO URV, O EL VALOR D'UNA EXPERIÈNCIA SOBRE LA JUBILACIÓ: COMPARTINT EMOCIONS

Seguim en contacte

El final d'una etapa laboral no sempre significa un comiat definitiu. Rosa Queral, coordinadora de Iubilo URV, explica els valors principals d'aquest projecte: garantir que l'experiència i la saviesa acumulades no es perdin, sinó que continuïn enriquint la comunitat universitària i reverteixin a la societat.

El meu últim dia de feina

L'any 2014 va arribar el dia que el meu lligam contractual amb la Universitat Rovira i Virgili va acabar. L'arribada de l'edat legal de jubilació amb els anys de cotització assolits és una gran oportunitat per continuar gaudint de la vida. La jubilació, coneguda com el *retiro obrero* (1919) és, juntament amb el dret als horaris racionals i el dret a les vacances, una de les grans fites assolides pels treballadors i les treballadores.

Havia començat com a professora titular de l'Escola d'Infermeria i, després d'estudiar Filosofia i Ciències de l'Educació i de fer el doctorat en Pedagogia, vaig acabar la meua vida acadèmica al Departament de Pedagogia. Amb la jubilació, el món dels sentiments amb la URV es va activar. Quan va acabar la meua relació contractual, la meua relació emocional amb la URV es va fer més viva i positiva.

L'últim dia, mentre arreplegava els llibres, els treballs, les fotos que m'havien acompanyat en el meu dia a dia i les anava col·locant en les capsas lliurades per companyes de la copisteria —companyes que sempre havien tingut un somriure i una paraula amable amb mi— vaig recordar que, abans d'ocupar el despatx per primera vegada, havia tingut la necessitat de renovar-lo. Un mes d'agost, amb la complicitat de la netejadora, en vaig pintar les parets de blanc, vam treure anys de pols i els records de la silueta dels marcs dels quadres de l'usuari anterior, hi vaig fer entrar la llum i l'alegria.

I feia poc que l'informàtic m'havia advertit que, al cap d'uns dies, l'adreça del meu correu electrònic desapareixeria i que perdria la informació que no tingués adesada en un llapis de memòria. El mateix tècnic que, amb paciència, m'havia anat ensenyant els programaris nous i m'havia anat aclarint dubtes sobre el funcionament de l'ordinador que se m'havien plantejat al llarg de la vida professional, m'informava que de cop i volta, d'un dia per l'altre, es tallaria el compte de correu electrònic que, durant molts anys, m'havia connectat amb el món.

En definitiva, el dia de la meva jubilació informativa de la URV vaig tancar el despatx no sense acomiadar-me de les parets, dels mobles i dels enormes xops o pollancres que, absorta en mil pensaments, contemplava des de la finestra.

Vaig tancar el despatx, en vaig lliurar les claus a les meves estimades companyes de la Secretaria i em vaig acomiadar dels companys i companyes del Departament amb abraçades i desig de bons auguris. Així acabava la meva vida laboral.

Aleshores va començar la vida de jubilada: una etapa en què aprenentatges nous van portar benestar i en què vaig veure que posar vida als anys era una fita a assolir.

La desconexió

El cert és, però, que jo no comptava amb la desconexió que causa la jubilació. Em vaig trobar en situacions tan paradoxals com no

assabentar-me de conferències interessants del meu departament ni de la meva facultat, ni de lliçons magistrals plenes de saviesa dels honoris causa, ni d'aquells esdeveniments universitaris que enforteixen les facultats mentals i psíquiques de les persones; facultats que, tot i ser vella, jo no havia perdut.

El privilegi d'estar viva

La utilització de la capacitat intel·lectual de totes les persones que conformen una societat és una oportunitat perquè aquesta es desenvolupi de forma sana i pròspera. Ancestralment, aquesta oportunitat ha estat negada a les dones —encara ho està en molts països teocràtics—, i avui en dia encara costa que sigui clarament reconeguda a les persones ancianes. Els estereotips i els prejudicis socials volen provocar un cert complex de culpa a la gent vella.

Efectivament, tant els homes com les dones ancians s'abstenen sovint d'expressar el potencial que tenen, afectats d'una estranya vergonya. Una vergonya que algú els provoca malintencionadament per apartar-los del món, com si el talent de les persones majors pogués fer alguna nosa a la civilització; com si l'expressió de l'experiència acumulada pels més grans pogués arribar a dificultar que els joves fessin seu el futur que els pertany.

La vellesa, un do

«No, la vellesa no és una malaltia, una incapacitat ni quelcom indigne que hagi d'averkonyir. Al contrari: la longevitat és un privilegi, un do que cal acceptar i aprofitar».

El 1986, la Premi Nobel de Medicina Rita Levi-Montalcini va afirmar amb rotunditat que les facultats del cervell humà són, inclús en una edat avançada, molt superiors a les que se li reconeixen, i que és necessari que el cervell no es jubili mai. Els nous coneixements sobre les connexions neuronals obren un camí d'esperança per trencar la idea que el jubilat o la jubilada és només aquella per-

sona retirada que espera pacientment l'arribada de la discapacitat, la dependència i la mort. No, la vellesa no és una malaltia, una incapacitat ni quelcom indigne que hagi d'avergonyar les persones que hi arriben. Ans al contrari: la longevitat és un privilegi que no es concedeix a tot el món: és un do que cal acceptar i aprofitar.

La motivació per a la fundació de Iubilo URV

En aquell moment em vaig fer conscient que, per gaudir del privilegi de la longevitat, és necessari fomentar la creativitat i tenir el convenciment que la lluita per aconseguir un món amb millors condicions per a tothom és possible. Vaig pensar que una institució amable com ho és la Universitat pot ser un bon entorn per continuar aquesta lluita, i vaig intuir que els seus jubilats i jubilades possiblement tenien les mateixes inquietuds que jo. Amb aquests pensaments, vaig guiar les meves passes cap a alguns membres de la URV.

Un bon dia de l'any 2018, havent expressat les meves inquietuds a una bona amiga que participava en la candidatura al Rectorat, ella em va dir: «Escriu tot això que penses i dius». El tríptic de campanya electoral va recollir de manera sintètica la meva proposta de crear Iubilo, nom que significa 'alegria'.

Em vaig alegrar d'haver pogut arribar a assolir una nova etapa vital, i d'haver sigut capaç de crear nous espais de connexió.

Creació de Iubilo URV

Una de les característiques del nostre projecte era d'obrir la possibilitat de ser membres de Iubilo URV a tot el personal procedent del PDI i del PAS que manifestés la voluntat de voler ser-ho. Fins aleshores, l'única manera de poder continuar la vinculació amb la URV, de gaudir d'un correu electrònic i accés a la biblioteca, havia estat la de demanar ser «Amiga o Amic de la URV», solució que,

després del vincle tan intens generat per anys de treball, semblava més una broma que una oportunitat.

La segona necessitat expressada a la vicerectora va ser la de poder mantenir el compte de correu electrònic de tota la vida, com passava en altres universitats, per exemple la UB, i que se suprimís la carta freda i poc amable que el Servei d'Informàtica enviava ritualment als seus usuaris a punt de jubilar-se per advertir-los de la cancel·lació.

També es va demanar que es fes un acte especial d'acomiadament dels companys i companyes en el moment que es jubileessin, que fos més expressiu que la simple i comprimida menció als jubilats i a les jubilades a les inauguracions del curs acadèmic.

Tot recollint les nostres peticions, en la sessió 10/2018 del Consell de Govern, que es va fer el 22 de novembre del 2018, l'equip rectoral de la Universitat Rovira i Virgili va crear la comunitat Iubilo URV amb l'objectiu de millorar el benestar de la comunitat universitària.

Per la raó d'haver estat jo qui havia proposat la creació de Iubilo URV, la rectora Figueras em va donar l'oportunitat d'impulsar el projecte. A tal fi, em va nomenar professora *ad honorem* i, juntament amb el Jaume Aymí, ens encomanava de dur a terme el projecte Iubilo URV.

Primers passos i desenvolupament de Iubilo URV

La primera gestió que vam emprendre l'any 2018 va ser la d'oferir als jubilats i jubilades de la URV la possibilitat d'incorporar-se a Iubilo URV. En l'acte públic de reconeixement que se'ls va fer el 7 de novembre del 2018, la Dra. Rosa Sánchez Rodríguez, de la Facultat de Ciències Jurídiques, s'hi va referir expressament, i el *Diari Digital de la URV* se'n va fer ressò tot anunciant que Iubilo URV seria coordinat pels professors jubilats Jaume Aymí i Rosa Queralt amb una comissió gestora d'entre sis i vuit PDI i PAS jubilats.

Seguidament, es va celebrar una reunió oberta a totes aquelles persones que es volien afegir al projecte Iubilo URV, prèviament convocada per la vicerectora encarregada de l'Oficina de Compromís Social. Les jubilades que s'hi van sumar i es van oferir voluntàries per formar la Junta executiva de Iubilo URV serien Àngels Olivé, Encarna Sastre, Esther Forgas, M. Isabel Miró i Liz Russell.

El 14 de maig del 2019, en ocasió de visitar «Girona, Temps de Flors», vam contactar amb les companyes i companys jubilats de la Universitat de Girona, que ens van fer conèixer la seva organització: una associació sense ànim de lucre anomenada Consell d'Amics i Amigues Seguint Fent UdG. Nosaltres vam preferir, però, mantenir el personal jubilat dins la comunitat universitària, amb totes les conseqüències que això comporta.

Els nostres suggeriments es van tenir en compte en el reglament, que es va aprovar el 18 de juny del 2019, i en el projecte d'Estatut que acabaria sent aprovat pel Claustre Universitari el 4 de novembre del 2021 i que es publicà al DOGC núm. 8623-10.3.2022.

Les referències que l'Estatut fa a Iubilo URV són al «Títol V. Comunitat universitària», al «Capítol 1. Disposicions generals». Són les següents:

Article 107. *Definició*

1. La comunitat universitària està formada per tres col·lectius: l'estudiantat, el personal docent i investigador i el personal d'administració i serveis.
2. Les comunitats *Alumni* i Iubilo es consideren vinculades a la Universitat d'acord amb els programes aprovats pel Consell de Govern.

Article 109. *Vinculació dels membres dels col·lectius Alumni i Iubilo.*
Al punt 2 diu: El personal jubilat de la Universitat que manifesti la voluntat de continuar-hi vinculat formarà part de la comunitat Iubilo, d'acord amb la normativa. D'aquesta manera, continuarà sentint-se part de la Universitat, i la institució continuarà comptant amb l'expertesa, saviesa i professionalitat de qui n'ha format part.

Finalitats de Iubilo URV

Em refereixo seguidament als objectius de la Comunitat Iubilo URV:

- 1) **Mantenir la connexió.** Segons la Dra. Rita Levi-Montalcini, «quan estimes allò que fas, mai tens la sensació que estàs treballant». Moltes persones que hem treballat a la nostra universitat ens hem sentit afortunades de poder dedicar-nos al servei de la societat. Creiem que el nostre coneixement no és un tresor amagat per gaudir-ne nosaltres mateixos, i que la nostra experiència i coneixement no es pot limitar a la nostra individualitat, perquè només quan és compartit podem aprofitar-ne el vertader potencial i generar un impacte positiu a la nostra societat. En compartir-lo no solament ajudem les noves generacions a créixer i desenvolupar-se, sinó que ens ajudem a nosaltres mateixos i contribuïm al progrés individual i col·lectiu.
- 2) **Fomentar un sentit ampli de la comunitat URV a la societat.** Els jubilats i jubilades que formem la comunitat Iubilo URV som un conjunt de persones grans amb necessitats, riscos i altres condicions diverses, però amb un factor que ens uneix: la vida a la URV. Som persones grans amb necessitats diferents que, amb esforç, hem assolit al llarg de la nostra vida experiències i sabers, i que hem fet possible que la nostra universitat sigui com és i on és. Les persones que l'hem fet possible volem que la URV sigui un motiu d'orgull per a les noves generacions. Els qui, des de la URV, hem estat al servei del coneixement i la millora de la societat mereixem continuar vinculades a la comunitat universitària.
- 3) **Reconèixer la capacitat del professorat i del personal d'administració i serveis jubilat d'oferir serveis a la comunitat universitària.** Es reté el talent de les persones que, altrament, es desvincularien del món acadèmic després de jubilar-se.

La societat actual, altament dinàmica gràcies a la velocitat vertiginosa del desenvolupament científic i tecnològic, podria arribar a marginar la gent gran per raó dels seus eventuais dèficits en coneixements tecnològics. Aquest malentès menystindria les experiències útils que aquestes persones podrien aportar a les noves generacions. L'impuls de la recerca interdisciplinària i multidisciplinària és important.

El cas és que, fins ara, la difusió de l'experiència del nostre col·lectiu de professorat i PAS universitari jubilat més aviat pot animar les noves generacions a impulsar la recerca per a la creació de productes, serveis i pràctiques innovadores per a l'atenció a la vellesa en àmbits com la salut, l'enginyeria, l'arquitectura, les ciències socials, etc.

- 4) **Promoure la jubilació activa organitzant activitats per a la comunitat universitària i en l'àmbit d'influència de la URV.** Les esferes del coneixement que ens interessen permeten actualitzar-nos en les novetats, col·laborar en la transferència de coneixement, en la formació dels professionals de la salut, de l'educació, de l'arquitectura, de l'enginyeria... que han de prestar serveis en l'atenció diària al col·lectiu de les persones grans, etc. Altrament, la jubilació activa és un concepte administratiu, però a Iubilo URV es tradueix en la necessitat de conèixer l'aplicació de la recerca científica a les persones grans i les seves famílies i els avenços i evidències dels àmbits professionals que poden estar a la seva disposició. Pensem que Iubilo URV pot ser un focus d'informació dels avantatges i una alerta dels possibles inconvenients administratius de l'etapa de jubilació.
- 5) **Mantenir present l'experiència de servei del personal jubilat de la URV,** en forma de col·laboració amb els òrgans de govern, facultats, instituts, serveis... per mantenir el bon nivell adquirit per la URV. El fet d'animar les persones grans a mantenir la bona salut física, psíquica i social implica la mo-

bilització de tots els àmbits socials perquè es comprometin a adoptar les mesures i objectius que permetin assolir fites concretes. Una bona fórmula per neutralitzar la presumpció de «caducitat» amb què de vegades es tracta socialment la vellesa és buscar models positius que fomentin un estil de vida a través del qual una persona se senti viva mantenint els seus coneixements al dia. Cal treballar conjuntament per aconseguir una situació òptima d'acord amb les necessitats, desitjos i capacitats i proporcionar protecció, seguretat i cures adequades quan facin falta. Per aquesta raó, la comunitat Iubilo URV ha de proporcionar la col·laboració necessària per promoure la solidaritat intergeneracional, el suport i la cooperació mutus entre les persones de diferents edats, segons les seves necessitats i capacitats, i beneficiar-se del progrés econòmic i social en igualtat de condicions.

- 6) **Reactivar el coneixement generat pel col·lectiu de personal jubilat que forma part de la URV.** Per incrementar el valor científic i social de la URV en la societat, Iubilo URV aporta el valor de la vellesa, del benefici de la salut emocional i de la feina feta per la URV. És per aquesta raó que Iubilo URV va participar en la Lliga de Debat Sènior de la Xarxa Lluís Vives. Aquesta participació va dotar els companys i companyes que van formar l'equip d'una oportunitat de millora de les seves habilitats comunicatives i d'autorealització, va afavorir un intercanvi d'experiències entre el grup i la resta de grups participants, i va ser una ocasió per mostrar i recollir informació del col·lectiu Iubilo URV i, en definitiva, d'aplegar un bon grau d'autoestima, ben merescuda.
- 7) **Actuar com a espai de debat sobre temes d'interès socio-cultural, en general, i sobre temes d'interès universitari, en particular.** El coneixement i l'experiència són béns preuats. El professorat i el personal d'administració i serveis constitueixen una valuosíssima font d'informació i de possibilitats

d'assessorament en els processos d'innovació. Investigar, fomentar l'educació formal i no formal en què participen les persones grans, sigui a les ciutats o en àmbits rurals, en les Aules de Gent Gran, requereix estudiar les necessitats sentides, expressades, comparades i normatives referents a la gent vella. Això implica contactar amb els grups de recerca de la URV per esbrinar les diferents línies obertes, donar a conèixer els resultats obtinguts entre els col·lectius, detectar desviacions i proposar millores de recerca i animar a dur a terme enquestes de satisfacció entre els usuaris i usuàries dels serveis prestats per la URV. Caldria crear, com han fet altres universitats, els estudis sènior en cultura, ciència, tecnologia i societat. Uns estudis que no tenen objectius professionalitzadors, però donen la possibilitat d'obrir la Universitat a aquelles persones que no han pogut desenvolupar estudis universitaris per qüestions econòmiques, laborals, o per falta d'oportunitats (el col·lectiu femení en pot ser un cas evident). Per altra part, pot facilitar que gent amb estudis universitaris previs puguin aprofundir en altres disciplines diferents d'aquelles que han estudiat. La vertadera riquesa d'una institució com la URV ha de ser l'impacte positiu que és capaç de generar en la vida de les persones. Això implica reavaluar les prioritats i buscar que l'èxit de les nostres actuacions com a universitat pública vagi més enllà dels factors econòmics per construir un món més humà i solidari.

Transferència

La vertadera riquesa d'una institució com la URV ha de ser l'impacte positiu que és capaç de generar en la vida de les persones.

- 8) Ser un referent i un espai de trobada i lleure de les persones jubilades de la URV.** La solitud no desitjada és i ha estat un motiu de debat en els àmbits científics i socials. Restar al

marge de la societat, no tenir ningú amb qui compartir un projecte vital, un espai per confrontar idees, coneixements, pensaments, per sentir-se part d'una xarxa social o senzillament per gaudir de moments d'esbarjo, comporta situacions que generen angoixes, pors i factors de risc que tenen conseqüències negatives importants en la qualitat de vida i la salut. A la URV, Iubilo potencia la comunicació, perquè organitza activitats que fomenten el cara a cara, crea xarxa social i, tot i que les noves tecnologies de vegades són percebudes com a elements d'aïllament perquè trenquen la possibilitat d'aquest cara a cara, poden ajudar quan per l'estat de salut i la pèrdua d'autonomia o de mobilitat es necessiten com una bona eina per mantenir la connexió. Iubilo URV té la vocació de promocionar l'autoestima i evitar la solitud social de les persones jubilades de la URV. Donar veu al personal de la Universitat Rovira i Virgili jubilat perquè s'organitzi, dinamitzi activitats per al col·lectiu i desenvolupi un servei a la Universitat i a la societat, en la mesura de les seves possibilitats, no és res més que fomentar l'envelliment actiu, tant de manera individual com col·lectiva.

Les persones que formem la comunitat Iubilo URV

El nombre de persones adherides a Iubilo URV en data 19 d'octubre del 2024 és de 312, de les quals 149 són dones (48 %) i 163 són homes (52 %).

El PAS el representen un total de 50 persones (16 %) i el PDI, 262 (84 %).

Els llocs de residència de les persones que formen la comunitat Iubilo URV són els de la Taula 1.

L'accés del PDI i el PAS a Iubilo URV es fa a través del formulari de la web. Es va donar el cas que molta gent que no havia estat mai PDI o PAS de la nostra universitat o de la Divisió VII de la Universitat de Barcelona omplien el formulari per inscriure-s'hi.

Comunicar a les persones que havien omplert el formulari sense ser membres de la universitat que no podien formar-ne part ens va portar a idear els Amics i Amigues de Iubilo.

Aquest darrer any hem tornat a revisar el formulari per demanar la data de jubilació o de finalització de situacions d'emèrits o ad honorem, perquè les persones que preveïen jubilar-se, abans de la data, ja demanen l'ingrés a Iubilo URV, amb el consegüent inconvenient tant per a la gestió de Iubilo URV com per a l'Oficina de Compromís Social i Recursos Humans, ja que no es pot donar d'alta a ningú sense la condició efectiva de jubilat o jubilada. Aquesta situació portava, per exemple, a expedir carnets quan encara era vigent el de PDI o PAS, extraviar formularis, etc.

Taula 1. Poblacions de residència de les persones adherides a Iubilo URV, amb el nombre (N) a cadascuna, agrupades per comarques o àmbits territorials.

POBLACIÓ	N	POBLACIÓ	N	POBLACIÓ	N
Tarragona	189	La Pobla de Montornès	2	Renau	1
Altafulla	8	El Catllar	1	Torredembarra	1
Salou	5	La Secuita	1	Vila-seca	1
Els Pallaresos	4	Perafort	1	Vilallonga del Camp	1
Creixell	3				
Reus	37	Mont-roig del Camp	2	L'Hospitalet de l'Infant	1
Cambrils	7	Riudoms	2	Vilaplana	1
Castellvell del Camp	4	L'Albiol	1	Vinyols i els Arcs	1
Les Borges del Camp	3	L'Aleixar	1		
La Selva del Camp	2				
Tortosa	7	Valls	4	Vilafranca del Penedès	2
Jesús	1	Alcover	2	Les Cabanyes	1
Roquetes	1	El Vendrell	3	Sant Martí Sarroca	1
Ulldecona	1	Falset	2	Sant Pere de Riudebitlles	1
		Bellmunt del Priorat	1	Vilanova i la Geltrú	1
Barcelona	12	Vilaverd	1	Igualada	1
Sant Cugat del Vallès	1				
Calella	1	Girona	1		
Rupit i Pruit	1	Olot	1	Santander	1
		Fortià	1		

Sobre els Amics i Amigues de Iubilo URV

Els Amics i Amigues poden participar en els actes oberts, i en el cas de no cobrir places en alguns esdeveniments, les oferim a aquest col·lectiu. En alguns casos en què no hem cobert les suficients places, han estat de gran ajuda. El nombre d'Amics i Amigues a data 19 d'octubre del 2024 és de 47 persones, de les quals 8 són homes (17 %) i 39, dones (83 %).

Sobre les activitats de Iubilo URV

Les activitats lúdiques són un bon antídote per combatre la solitud no volguda, aproximar persones, establir lligams i, sobretot, gaudir de la joia de viure.

Participar en jocs i debats ens ajuda a aprendre. El fet d'aprendre és un procés que es manté mentre la vida dura i cal estimular-lo.

La vida contemplativa és tan real com la vida activa i és una necessitat universal que alimenta l'esperit i l'intel·lecte. Les visites culturals tenen el seu fonament en aquesta necessitat universal.

La carta més valuosa per a les persones està representada per la capacitat de valdre'ns, psíquicament i mentalment, en totes les etapes i, en especial, quan som ancians i ancianes. El testimoni vital de moltes persones conegudes en tots els camps del saber ens demostra que la capacitat cognitiva perdura encara que amb formes noves en edats avançades. I és aquí on rau l'esperança i la fortalesa de Iubilo URV, la de poder posar les cartes sobre la taula i lluitar perquè la intel·ligència, la paciència, l'eloqüència, l'estima, la noblesa, la intuïció, l'alegria, més aviat o més tard, ens ajuden a intuir una societat de persones lliures.

10. LA XARXA VIVES D'UNIVERSITATS I LA LLIGA DE DEBAT SÈNIOR

Punt de trobada

La Xarxa Vives d'Universitats, creada el 1994, coordina l'acció conjunta de 22 universitats de territoris de parla catalana. Amb seu a la Universitat Jaume I de Castelló, fomenta la cooperació acadèmica i la normalització lingüística, i organitza activitats com les Lligues de Debat.

La Xarxa Vives d'Universitats és una institució sense ànim de lucre que avui dia representa i coordina l'acció conjunta de 22 universitats.

Es va constituir el 28 d'octubre del 1994 a Morella, en aquell moment amb 13 rectors d'universitats de parla catalana com a socis fundadors: entre altres, la nostra URV, amb el rector de llavors, Joan Martí; Pep Nadal per la UdG; Carles Solà per la UAB i Antoni Caparrós per la UB. En un primer moment es va anomenar Institut Lluís Vives, en honor de l'humanista valencià, i al llarg dels anys s'ha convertit en la Xarxa Vives. Té la seu organitzativa a la Universitat Jaume I de Castelló.

És una plataforma que aplega institucions universitàries de Catalunya, el País Valencià, les Illes Balears, la Catalunya del Nord, Andorra i Sardenya, i també d'altres territoris amb vincles geogràfics, històrics, culturals i lingüístics comuns, per tal de crear un espai universitari que permeti coordinar la docència, la recerca i les activitats culturals i potenciar la utilització i la normalització de la llengua pròpia.

La raó fonamental que va guiar la creació d'aquest organisme va ser la necessitat de coordinació d'activitats en tot aquest àmbit d'actuació, amb un espai telemàtic comú i amb la intenció d'organitzar un servei de publicacions destinades a la docència superior.

La Xarxa Vives d'Universitats desenvolupa la seva missió:

- Mitjançant una sistemàtica de treball coordinat i compartit, fonamentat en l'equitat entre els seus membres, la qualitat i el servei.
- Formulant propostes conjuntes de les universitats i facilitant la seva inclusió en l'agenda social i política de cada realitat administrativa.
- Contribuint a fer encara més forta l'educació superior i la recerca.

Els seus valors són:

- Unitat lingüística de les comunitats universitàries membres, amb vincles geogràfics, històrics i culturals comuns.
- Col·laboració interuniversitària.
- Equitat entre totes les universitats.
- Respecte a l'autonomia universitària.
- Democràcia com a forma de govern de la institució.
- Voluntat de projecció i servei conjunt a la societat.

La unitat lingüística és el principal valor fundacional sobre el qual es va crear la Xarxa, i ja des de l'inici es dota d'una Comissió de Llengua que ha desenvolupat de manera continuada la seva tasca de coordinació, intercanvi, formació i projecció de la llengua catalana.

La Xarxa Vives ha estat pionera a desenvolupar un pla de política lingüística propi, consensuat per 22 universitats d'un mateix domini lingüístic de quatre estats, que incorpora criteris generals en matèria de llengua, aplicables a tots els àmbits de la universitat, i que vetllen per la normalització, el multilingüisme, la coordinació d'accions i la promoció del català a través de més de 60 línies de treball cada any. També s'hi impulsen accions per potenciar les

relacions entre les principals entitats lingüístiques i la interlocució i mediació entre elles, les universitats i les entitats governamentals en matèria de política lingüística.

Impuls al català

La Xarxa ha desenvolupat un pla de política lingüística propi per promoure el català en l'àmbit acadèmic i fomentar el multilingüisme amb accions conjuntes entre universitats i institucions.

Entre moltes altres activitats dutes a terme en aquests anys, destaca especialment el treball amb els estudiants, amb la Lliga de Debat de Secundària i de Batxillerat (amb edicions des de l'any 2010) i la Lliga de Debat Universitària (amb edicions des de l'any 2005). Pel que fa al nostre àmbit sènior, el que ens ha acabat d'integrar realment a la Xarxa ha estat la Lliga de Debat Sènior, que enguany acaba de celebrar la tercera edició, i que tot seguit comentarem.

Els Debats a la Xarxa Vives: què són?

Les Lligues de Debat són les competicions dialèctiques entre equips d'estudiants que la Xarxa Vives d'Universitats organitza sobre temes polèmics i d'actualitat amb la finalitat de fomentar l'ús de la paraula, potenciar el treball en equip, estimular l'argumentació i la posada en escena i afavorir la relació entre els estudiants. La Xarxa convoca anualment diferents lligues, que adreça, respectivament, a estudiants de 4t d'ESO, batxillerat i formació professional, a estudiants universitaris i a estudiants de programes sènior.

La Lliga de Debat de Secundària i Batxillerat és una competició d'oratori en la qual diversos equips d'estudiants de 4t d'ESO, batxillerat i formació professional de diferents centres docents debaten entre si sobre un tema polèmic i d'actualitat. L'objectiu és de fomentar les competències transversals dels estudiants mitjançant l'ús de la paraula, i de crear un espai de formació on s'utilitza el català com a llengua vehicular.

Aquest format té ja una bona tradició des de l'any 2010. Els estudiants que hi han participat tot representant la URV hi han guanyat diversos premis: l'any 2019 l'Institut de Tecnificació d'Amposta va ser subcampió, i l'any 2022 Jaume Alexandre Custodi, de l'Institut Serra de Miramar de Valls, va ser el millor orador.

La Lliga de Debat Universitària és una competició dialèctica en la qual, durant una setmana, diferents equips d'estudiants universitaris debaten sobre un tema polèmic i d'actualitat. L'objectiu de la competició és fomentar l'ús de la paraula i potenciar capacitats com el treball en equip, l'argumentació i la posada en escena i la relació amb estudiants d'altres universitats.

La Xarxa Vives ha celebrat debats al món universitari des de l'any 2005. La URV va guanyar l'edició de l'any 2017, en què es va debatre el tema «Cal que el poder públic controli els mitjans de comunicació?».

La Lliga de Debat Sènior és una competició dialèctica en la qual diversos equips d'estudiants de programes sènior de diferents universitats de la Xarxa Vives debaten entre ells sobre un tema polèmic i d'actualitat.

Des d'un enfocament basat en l'envelliment actiu, reconeix que les necessitats intel·lectuals, culturals i socials no minven amb el pas del temps i que activitats formatives com aquesta poden cobrir perfectament aquestes necessitats de les persones grans i acomplir una funció social inclusiva i de ciutadania plena.

Aquesta competició anual afavoreix l'intercanvi de bones pràctiques en l'atenció a les persones majors, i dona suport als programes de promoció de l'aprenentatge permanent que ofereix, concretament, cadascuna de les universitats associades.

En aquest sentit, cada universitat dissenya el seu propi programa sènior en funció de les necessitats que vol cobrir, però totes elles coincideixen en la intenció d'ajudar les persones grans a millorar el seu desenvolupament personal i augmentar la seva qualitat de vida.

Vistos en conjunt, els programes sènior de la Xarxa mostren l'interès de la comunitat per incorporar, mantenir o reincorporar al món universitari el talent de les persones que ja estan a punt d'acabar la vida laboral o que ja en són fora.

En aquest sentit, a la URV ja estem participant en algunes activitats del programa, com són les Aules de la Gent Gran, però ens cal encara una major participació en forma de millor estructuració del nostre sistema d'alumnes sènior (URV Ciutadana).

Quins són els objectius de la Lliga de Debat Sènior?

Els objectius són millorar les habilitats comunicatives dels participants, fomentar l'adquisició de competències transversals, afavorir la convivència i l'intercanvi d'experiències entre alumnat sènior i, en definitiva, oferir als estudiants sènior oportunitats de creixement i d'autorealització que en milloren la qualitat de vida.

Qui pot participar en la Lliga de Debat Sènior?

La participació en la competició està oberta a l'alumnat dels programes sènior de les universitats membres de la Xarxa Vives que estigui matriculat en el curs acadèmic en què es fa la competició.

Els equips representen la seva universitat i han d'estar formats per un mínim de dues i un màxim de cinc persones. Aquest equip haurà de comptar amb un capità o capitana, que en farà de portaveu/representant.

Com es desenvolupa la Lliga de Debat Sènior?

Uns mesos abans del començament de la Lliga es fa públic el tema del debat. Durant la competició, els equips no coneixen quina posició hauran de defensar (sigui a favor o bé en contra) fins a uns minuts abans de començar cada enfrontament. Un equip de jutges decideix quina argumentació s'ha fet amb el màxim rigor i de la forma més convincent possible, d'acord amb els criteris establerts en les bases de la competició.

La primera edició es va celebrar l'any 2022 a la Universitat Autònoma de Barcelona (UAB). El tema de debat va ser «És necessària una indústria farmacèutica pública a la Unió Europea?». L'equip campió va ser el de la Universitat Miguel Hernández, d'Elx.

La segona edició va tenir lloc l'any 2023 a la Universitat Miguel Hernández d'Elx, seu de l'equip campió de l'edició anterior. El tema a debatre va ser «L'humor té límits?». L'equip campió va ser el de la Universitat de Vic-Universitat Central de Catalunya.

11. BREU GLOSSARI DE TERMES DE LA IA

Això és un glossari de termes genèrics d'intel·ligència artificial (IA), amb els enfocaments d'IA més importants, ordenats alfabèticament. S'hi inclouen els noms més comuns en anglès en cursiva, i són subratllats els termes de la llista esmentats en altres definicions.

Algorisme o algoritme: És el conjunt finit d'instruccions o passos per resoldre un problema específic o processar càlculs o dades.

Aprenentatge automàtic (*machine learning*): Camp de la IA per al disseny, anàlisi i desenvolupament d'algorismes que permeten que els programes dels maquinaris evolucionin estadísticament.

Aprenentatge profund (*deep machine learning*): Tècnica d'extracció i transformació d'aprenentatge automàtic, per mitjà d'algorismes per capes, que simula el funcionament neuronal. Va ser clau (2012) en l'inici del boom de la IA pocs anys després.

Assistent virtual: Bot (abreviatura de *robot*) digital de conversa, sobretot els més clàssics, previs als models de llenguatge extens, com ELIZA (1964 al MIT), Cortana (Microsoft), Assistant (Google), Alexa (Amazon), Celia (Huawei), Siri (Apple), Tobi (Vodafone) i altres.

Dades massives (*big data*): Conjunts de dades, procediments i aplicacions informàtiques de gran volum, que ultrapassen la capacitat dels sistemes habituals, i requereixen eines d'IA.

IA feble o IA estreta (*artificial narrow intelligence*): IA creada per fer unes tasques concretes amb una alta eficiència. Ho són totes les IA conegudes actualment.

IA general o IA forta, (*artificial general intelligence*): IA que comprendria i portaria a terme la major part de tasques intel·lectuals que pot fer un humà, de forma general. Per ara, només existeix a la ciència-ficció.

IA generativa: IA capaç de generar textos, imatges, vídeos o altres dades utilitzant models generatius, com els models de llenguatge extens. Cal no confondre-la amb la IA general.

IA no simbòlica o **connexionista**: IA basada en l'aprenentatge automàtic (sense una programació explícita) i el processament del llenguatge natural. Es fa amb algorismes que identifiquen patrons a les dades. Ha tingut un desenvolupament ràpid els darrers anys, sobretot la IA generativa.

IA simbòlica: IA basada en el raonament lògic i la representació del coneixement. Ha estat la base de molts sistemes d'èxit, com els sistemes experts utilitzats en biologia molecular, medicina i finances.

Mineria de dades (*data mining*): Procés d'extreure i descobrir patrons en grans conjunts de dades, utilitzant aprenentatge automàtic, estadística i sistemes d'organització de bases de dades.

Model de llenguatge extens (*large language model, LLM*): Model de llenguatge d'una xarxa neuronal amb molts paràmetres entrenats en grans quantitats de text mitjançant aprenentatge no supervisat. Per ex. la sèrie GPT-3, GPT-4 usats a ChatGPT.

Processament del llenguatge natural: Disciplina informàtica que tracta les llengües naturals.

Prova de Turing: Alan Turing (1912-1954) ho va establir així: Un avaluador fa preguntes a una altra persona i a un ordinador dissenyat per generar respostes anàlogues a les dels humans. L'avaluador no sap qui és la persona i qui l'ordinador, es comunica a través d'una interfície intermèdia. Si l'avaluador no pot distingir la persona de la màquina o considera que la màquina

és una persona, la màquina superaria el test de Turing i es consideraria intel·ligent.

Sistema expert: Sistema informàtic —conjunt de programes— que emula el comportament d'un expert en una àrea específica, com diagnòstic mèdiques, planificació de finances, química industrial, etc.

Súper IA (*artificial super intelligence*): Hipotètica IA amb capacitats superiors a la intel·ligència humana (Mucci & Stryker, 2023).

Transformadors generatius preentrenats (*generative pre-trained transformer, GPT*): Família de models de llenguatge extens (LLM) entrenats en un corpus de dades de text per generar text semblant a l'humà.

Xarxa neuronal o XN artificial (*neural network, NN*): Model matemàtic d'aprenentatge automàtic (ML) inspirat en el sistema nerviós biològic, amb nòduls interconnectats.

Bot de conversa (*chatbot*) o assistent virtual intel·ligent: Programa, normalment en línia, que simula una conversa com a interlocutor amb un usuari en llenguatge natural, que pot ser per teclat o amb interfície d'usuari multimèdia. Aquests sistemes experts utilitzen el raonament basat en casos.

12. PREGUNTES A L'APLICACIÓ D'IA CONVERSACIONAL CHATGPT V2

Preguntes fetes a la versió lliure bàsica en català [ChatGPT v2](#) relacionades amb el capítol 6: «Conclusió: la IA és una amenaça o no?».

1. Sol·licitud de l'Albert feta l'11 de novembre del 2024

Fes un text d'unes 300 paraules que respongui a la pregunta següent: La intel·ligència artificial és una amenaça per a la humanitat?

Resposta del ChatGPT

La intel·ligència artificial (IA) ha generat un debat intens sobre el seu impacte en la humanitat, plantejant tant oportunitats com riscos. Des d'una perspectiva optimista, la IA representa una eina poderosa capaç de millorar la vida de les persones. Pot contribuir a avanços en camps com la medicina, amb diagnòstics més precisos i tractaments personalitzats; en la sostenibilitat, optimitzant recursos per combatre el canvi climàtic, o en l'educació, amb sistemes d'aprenentatge adaptatius que responen a les necessitats individuals. Tanmateix, aquesta mateixa capacitat per transformar sectors pot comportar riscos i, potencialment, esdevenir una amenaça per a la humanitat si no es gestiona de forma adequada.

Un dels perills més destacats és l'automatització excessiva que podria portar a la substitució de llocs de treball, amb repercussions importants en l'economia i la cohesió social. A més, l'ús de la IA en la vigilància massiva i el control de la població pot vulnerar la privacitat i les llibertats civils. També s'ha advertit sobre la possibilitat que els sistemes d'IA desenvolupin, de manera no intencionada,

comportaments perillosos o no alineats amb els interessos humans, sobretot en sistemes que funcionen de manera autònoma, com els vehicles o armaments.

La qüestió més preocupant és el desenvolupament d'una IA avançada que superi les capacitats humanes (superintel·ligència) i que, sense una regulació ètica i tècnica adequada, pugui tenir objectius diferents o incomprensibles per als humans, amb conseqüències imprevisibles.

En resum, la IA no és una amenaça per si mateixa, sinó que depèn de com la societat decideixi implementar-la i regular-la. Si s'inverteix en una recerca responsable i es desenvolupen normatives internacionals clares, la IA pot ser una gran aliada per a la humanitat. En canvi, si es desatén la seva gestió ètica i responsable, sí que podria comportar riscos significatius per a la nostra societat.

2) Sol·licitud de l'Albert feta el 20 d'octubre del 2024:

Fes un text d'unes 200 paraules màxim que comenti els principals avantatges de la intel·ligència artificial per als majors de 65 anys.

Resposta del ChatGPT

La intel·ligència artificial (IA) aporta nombrosos avantatges per a les persones majors de 65 anys, millorant la seva qualitat de vida i fomentant la seva autonomia. Un dels principals beneficis és en l'àmbit de la salut, on la IA pot ajudar a fer un seguiment mèdic personalitzat, recordant la presa de medicaments o controlant signes vitals. Això és especialment útil per a persones amb malalties cròniques o necessitats d'atenció constant.

A més, la IA ofereix solucions per combatre l'aïllament social gràcies als assistents virtuals que permeten mantenir converses o realitzar tasques quotidianes a través de comandes de veu. Això facilita la interacció amb la tecnologia i redueix la barrera digital, ajudant les persones grans a mantenir el contacte amb familiars i amics.

També millora la seguretat domèstica amb sistemes automatitzats que monitoritzen la casa i poden avisar en cas d'emergència, com una caiguda. A més, la IA pot oferir entreteniment personalitzat i activitats cognitives, fomentant l'estimulació mental i l'aprenentatge constant.

En conjunt, la IA proporciona una major independència, seguretat i benestar emocional, i facilita la vida diària de les persones majors.

13. BIBLIOGRAFIA



Aquí tenim un codi QR que porta a un
lloc web amb aquesta mateixa relació de Bibliografia
amb els hiperenllaços directes a les referències.

<https://drive.google.com/file/d/1FzlU-TMU-MJIG8QtFHiig-GuOG2dvrUau/view?usp=share_link>

Bibliografia, per ordre alfabètic d'autor, referenciada en el text.

- ABBAS, A. (2024). «The next generation of tiny AI: Quantum computing, neuromorphic chips, and beyond». A: *Unite.AI*. 8/3/2024. <<https://www.unite.ai/the-next-generation-of-tiny-ai-quantum-computing-neuromorphic-chips-and-beyond/>>
- ABRAMSON, J. Et al. (2024). «Accurate structure prediction of biomolecular interactions with AlphaFold 3». A: *Nature*. 10.1038. <<https://www.nature.com/articles/s41586-024-07487-w>>
- ALBERCA, C. (2023). «Novedades en IA en ofimática. Adictos al Trabajo». A: *Autentia*. 8/6/2023. <<https://adictosaltrabajo.com/2023/06/08/novedades-en-inteligencia-artificial-en-ofimatica/>>
- ANDERSEN, K. (2014). «Enthusiasts and skeptics debate Artificial Intelligence». A: *Vanity Fair*. 26/11/2014. <<https://www.vanityfair.com/news/tech/2014/11/artificial-intelligence-singularity-theory?srltid=AfmBOopFPiAgedzsFLMFqTUR55QUL8Io0BBdZnVJQMqDjIA5sokjIrOR>>
- ANGULO, M. J. et al. (2024). *Guia per a l'aplicació de la IA a l'acció docent de la UOC*. UOC-eLearning Innovation Center. <https://openaccess.uoc.edu/bitstream/10609/149762/1/CA_Guia%20IA%20PDC.pdf>
- ASIMOV, I. (1950). *I, Robot*. <<https://www.britannica.com/topic/I-Robot>>
- ALTCHER, A. (2024). «El antiguo CEO de Google tiene una solución muy simple ante la aterradora idea de que la IA tome el control». A: *Business Insider*. 26/5/2024. <<https://www.businessinsider.es/ex-ceo-google-tiene-solucion-ia-no-tome-control-1387834>>
- BAS-WOHLERT, C. (2024). «This new AI predicts your life. Then it predicts your death». A: *Science Alert*. 25/3/2024. <<https://www.sciencealert.com/this-new-ai-predicts-your-life-then-it-predicts-your-death>>
- BECK, U. (2006). *La sociedad del riesgo global*. Siglo XXI Editores.

- BERRONDO-OTERMIN et al. (2023). «Application of AI techniques to detect fake news: a Review». A: *Electronics* 12, 5041. <<https://www.mdpi.com/2079-9292/12/24/5041>>
- BHAGAVAD-GITA (sense data). Poema sànscrit contingut en el Mahabharata. Enciclopèdia.cat. <<https://www.enciclopedia.cat/gran-enciclopedia-catalana/bhagavad-gita>>
- BORRÀS, P. (2024) «La intel·ligència artificial avança i s'estén a tota la cadena agroalimentària». A: *El Triangle*. 23/03/2024. <<https://www.eltriangle.eu/2024/03/23/la-intelligencia-artificial-avanca-i-sesten-a-tota-la-cadena-agroalimentaria/>>
- BOYLE, A. (2020). «AI vs. humans? Microsoft's Eric Horvitz sees the future as AI-human partnership». A: *GeekWire, Bot or Not?*. 18/2/2020. <<https://www.geekwire.com/2020/ai-vs-humans-microsofts-eric-horvitz-sees-future-ai-human-partnership/>>
- CARBONELL, E. (2024). Entrevista a Librújula de Antonio Iturbe. 3/3/2024. <<https://librujula.publico.es/eudald-carbonell-la-inteligencia-artificial-es-el-descubrimiento-mas-importante-de-la-humanidad-despues-del-fuego/>>
- CAUDAI, C. Et al. (2021). «AI applications in functional genomics». A: *Comput Struct Biotech*. J 19, 5762-5790. <[https://www.csbj.org/article/S2001-0370\(21\)00431-1/fulltext](https://www.csbj.org/article/S2001-0370(21)00431-1/fulltext)>
- CHOMSKY, N. et al. (2023). «The false promise of ChatGPT». A: *The New York Times, Opinion*. 8/3/2023. <<https://www.nytimes.com/2023/03/08/opinion/noam-chomsky-chatgpt-ai.html>>
- DANIEL, W. (2024). «DeepMind co-founder Mustafa Suleyman warns AI is a 'fundamentally labor replacing' tool over the long term». A: *Fortune*. 17/1/2024. <<https://fortune.com/2024/01/17/mustafa-suleyman-deepmind-ai-a-i-labor-replacing-tool-over-the-long-term/>>
- DEGEURIN, M. (2024). «AI-generated nonsense is leaking into scientific journals». A: *Popular Science*. 19/3/2024. <<https://www.popsoci.com/technology/ai-generated-text-scientific-journals/>>

- DEL Bosque, M. (2024). «Alba Meijide, la gallega que marca el camino de la IA: ‘El fin de la humanidad y de todos los trabajos vende mucho’». A: *El Mundo, Noticias de Yo Dona*. 17/5/2024). <<https://www.elmundo.es/yodona/actualidad/2024/05/17/663a04a7e4d4d8000f8b458e.html>>
- DEL Castillo, C. (2024a). «¿Cuánto consumen los centros donde “piensa” la IA? Tanto como para necesitar una central nuclear». A: *El Diario.es*. 9/3/2024. <https://www.eldiario.es/tecnologia/consumen-centros-piensa-inteligencia-artificial-necesitar-central-nuclear_1_10982427.html>
- DEL Castillo, C. (2024b). «La IA se asoma a su propia catarsis punto-com: Muchos valores están inflados». A: *El Diario.es*. 1/5/2024. <https://www.eldiario.es/tecnologia/inteligencia-artificial-asoma-propia-catarsis-puntocom-valores-inflados_1_11328246.html>
- DOLGIN, E. (2023). «‘Remarkable’ AI tool designs mRNA vaccines that are more potent and stable». A: *Nature News*. 2/5/2023. <<https://www.nature.com/articles/d41586-023-01487-y>>
- ESCALA Blog (2024) «Ventajas y desventajas de la IA en empresas». Nexus Integra. <<https://escala.com/blog/ventajas-y-desventajas-de-la-inteligencia-artificial-en-empresas>>
- ESPINOZA, J. L. & Dong, L. T. (2020). «AI tools for refining lung cancer screening». A: *J Clin Med* 9, 3860. <<https://www.mdpi.com/2077-0383/9/12/3860>>
- ESTUPINYÀ, P. (2024). «Hinton, una voz de alerta ante la IA». *Cazador de cerebros*. A: RTVE.es 2/6/2024. <<https://www.rtve.es/play/videos/el-cazador-de-cerebros/02-06-24/16129904/>>
- GARLAND, A. (2014). «Ex Machina, film». <[https://ca.wikipedia.org/wiki/Ex_Machina_\(pel·lícula\)](https://ca.wikipedia.org/wiki/Ex_Machina_(pel·lícula))>
- GIL, P. (2024). «La IA y el cine: una relación en constante evolución». A: *IA.CienciaTV.online*. 12/3/2024. <<https://ia.cienciatv.online/la-inteligencia-artificial-y-el-cine-una-relacion-en-constante-evolucion/>>

- GOTTFREDSON, L. S. (1997). «Mainstream science on intelligence: an editorial with 52 signatories, history and bibliography». A: *Intelligence*. 1997, 13-23. <[https://doi.org/10.1016/S0160-2896\(97\)90011-8](https://doi.org/10.1016/S0160-2896(97)90011-8)>
- HARARI, Y. N. (2014). *Sàpiens; una breu història de la humanitat*. Grup 62. Traducció de la versió en anglès del 2014, original en hebreu 2011.
- HAWKING, S. Et al. (2014). «Stephen Hawking: Transcendence looks at the implications of AI-but are we taking AI seriously enough?». A: *Independent News Science*. 1/5/2014. <<https://www.independent.co.uk/news/science/stephen-hawking-transcendence-looks-at-the-implications-of-artificial-intelligence-but-are-we-taking-ai-seriously-enough-9313474.html>>
- HISTORY Skills (2024). «Trains, terror, and technological change: 19th century hysteria about railways». Classroom Year 9. <<https://www.historyskills.com/classroom/year-9/steam-train-fears/>>
- HSU, J. (2024). «AI chatbots use racist stereotypes even after anti-racism training». A: *NewScientist*. 7/3/2024. <<https://www.newscientist.com/article/2421067-ai-chatbots-use-racist-stereotypes-even-after-anti-racism-training>>
- HUANG, S. C. et al. (2022). «Deep neural network trained on gigapixel images improves lymph node metastasis detection in clinical Settings». A: *Nature Commun.* 13, 3347. <<https://www.nature.com/articles/s41467-022-30746-1>>
- HUIZINGA, J. (1972). *Homo ludens*. Alianza Editorial.
- INSTITUT de Seguretat Pública de Catalunya (2017). «Capacitació per a la planificació de l'autoprotecció en l'àmbit local». Generalitat de Catalunya. <https://ispc.gencat.cat/ca/publicacions/planificacio-autoproteccio/capacitacio_planificacio_autoproteccio_ambit_local/>
- ISRAEL, R. (2023). «¿Debemos tenerle miedo a la IA?». A: *Infobae*. Laciencia.tv ¿Hay que tener miedo a la IA? 4/5/2023. <<https://>>

- www.infobae.com/america/opinion/2023/05/04/debemos-tenerle-miedo-a-la-inteligencia-artificial/
- KACENA, M. A. et al. (2024). «The use of AI in writing scientific review articles». A: *Curr Osteoporosis Rep.* 22:115. <<https://link.springer.com/article/10.1007/s11914-023-00852-0>>
- KAPLAN, A.; Haenlein, M. (2019). «Siri, Siri, in my hand: Who's the fairest in the land? On the interpretations, illustrations, and implications of artificial intelligence». A: *Business Horizons.* 62, 15–25. <<https://doi.org/10.1016/j.bushor.2018.08.004>>
- KHALIFA, M.; Albadawy, M. (2024). «Using artificial intelligence in academic writing and research: an essential productivity tool». A: *Comp Meth and Programs in Biomedicine Update.* <<https://doi.org/10.1016/j.cmpbup.2024.100145>>
- KURZWEIL, R. (2024). *The Singularity is nearer. When we erge with AI.* Random House. <<https://www.libreriainglesa.com/products/599273-the-singularity-is-nearer.html>>
- KUTA, S. (2024). «La IA acaba de descifrar parte de un antiguo pergamino “ilegible”: esto es lo que dice». A: *National Geographic.* 2/9/2024. <<https://www.nationalgeographic.es/historia/2023/10/inteligencia-artificial-desvela-secretos-pompeya-ia-descifrar-antiguo-pergamino-ilegible>>
- LAM, P. (2024). «The impact of AI on the art world». A: *Forbes.* 2/2/2024. <<https://www.forbes.com/sites/forbesbusinesscouncil/2024/02/02/the-impact-of-artificial-intelligence-on-the-art-world/>>
- LAWLOR, P.; Chang, J. (2023). «The positive social impact of AI; harnessing AI for the greater good and to improve our world». A: *Qualcomm.* 20/11/2023. <<https://www.qualcomm.com/news/onq/2023/11/the-positive-social-impact-of-ai>>
- LE Calme, S. (2024). «L'éditeur de revues académiques Wiley ferme 19 revues scientifiques et retire plus 11000 articles». A: *Developez.com.* 30/5/2024. <<https://intelligence-artificielle.developez.com/actu/358438/>>

- LIMÓN, R. (2024). «La evolución de la IA que hace parecer increíblemente tonto a ChatGPT y tiene riesgos éticos». A: *El País*. Tecnología, 13/5/2024. <<https://elpais.com/tecnologia/2024-05-13/la-evolucion-de-la-ia-que-hace-parecer-increiblemente-tonto-a-chatgpt-y-tiene-riesgos-eticos.html>>
- MACGREGOR, K. (2024). «Nobel Prize scientists on AI, democracy and critical thinking». A: *University World News*. 8/3/2024. <<https://www.universityworldnews.com/post.php?story=20240308135103305>>
- MICRONS Explorer (2024). «A virtual observatory of the cortex». Princeton University. <<https://www.microns-explorer.org/cortical-mm3>>
- MITCHELL, M. (2024). A: *Inteligencia Artificial: Guía para seres pensantes*. Ed Capitán Swing.
- MOR, R. et al. (2019). «Gaia DR2 reveals a star formation burst in the disc 2-3 Gyr». A: *Astronomy and Astrophysics*. 624, L1. <<https://doi.org/10.1051/0004-6361/201935105>>
- MUCCI, T.; Stryker, C. (2023). «What is artificial superintelligence?» A: IBM Topics. 18/12/2023. <<https://www.ibm.com/think/topics/artificial-superintelligence>>
- NATURE Editorial (2024). «Why scientists trust AI too much -and what to do about it». A: *Nature*. 627, 243. <<https://pubmed.ncbi.nlm.nih.gov/38448525/>>
- NEWS & Events (2024). «Study shows that the way the brain learns is different from the way that AI systems learn». A: *News & Events*. University of Oxford. 3/1/2024. <<https://www.ox.ac.uk/news/2024-01-03-study-shows-way-brain-learns-different-way-artificial-intelligence-systems-learn>>
- NIAZIAN, M.; Niedbala, G. (2020). «Machine Learning for Plant Breeding and Biotechnology». A: *Agriculture*. 10, 436. <<https://doi.org/10.3390/agriculture10100436>>
- NIELD, D. (2024). «AI detects mysterious detail hidden in famous Raphael masterpiece». A: *ScienceAlert*. 19/4/2024. <<https://>>

- www.sciencealert.com/ai-detects-mysterious-detail-hidden-in-famous-raphael-masterpiece>
- NIH (2024). «AI, Machine Learning and Genomics». A: *National Human Genome Res Inst.* NIH, US. <<https://www.genome.gov/about-genomics/educational-resources/factsheets/artificial-intelligence-machine-learning-and-genomics>>
- NOBEL Prizes (2024). «The Nobel Prize. Nobel Prizes 2024 in Physics and in Chemistry». <<https://www.nobelprize.org/all-nobel-prizes-2024/>>
- ONU (2015). «Objetivos de desarrollo sostenible. 17 objetivos para transformar nuestro mundo». *Agenda 2030*. <<https://www.un.org/sustainabledevelopment/es/objetivos-de-desarrollo-sostenible/#>>
- ORTIZ DE ZÁRATE, L.; GUEVARA, A. (2021). «IA e igualdad de género. Un análisis comparado entre UE, Suecia y España». A: *Estudios de Progreso, Fundación Alternativas*. 101. <<https://fundacionalternativas.org/wp-content/uploads/2022/07/cd41c-c86bb79705300ef0668114f037f.pdf>>
- PATRIZIO, A. (2024). «10 top AI jobs in 2024». A: *WhatIs? TechTarget*. 25/1/2024. <<https://www.techtarget.com/whatis/feature/Top-AI-jobs>>
- PÉREZ COLOMÉ, J. (2024). «Kyle Chayka, periodista: ‘En las industrias culturales no importa lo que haces, sino cuántos seguidores tienes’». A: *El País*. Tecnología. 6/5/2024. <<https://elpais.com/tecnologia/2024-05-06/kyle-chayka-periodista-en-las-industrias-culturales-no-importa-lo-que-haces-sino-cuantos-seguidores-tienes.html>>
- POMPAS, J. (2023). «¿La IA es realmente inteligente?» A: *Linkedin*. 25/8/2023. <<https://www.linkedin.com/pulse/la-inteligencia-artificial-es-realmente-inteligente-pompas-gutiérrez/>>
- PONCE, J. L. (2024). «¿Cómo aporta la inteligencia artificial a la ciberseguridad?» A: *Ikusi Velatia*. <<https://www.ikusi.com/es/blog/como-aporta-la-inteligencia-artificial-a-la-ciberseguridad/>>

- ROMERO-GÓMEZ, M. et al. (2019). «Gaia - La revolución del Big Data llega a la galàxia». A: *Astronomía*. SEA 244, 24-30. <https://www.sea-astronomia.es/sites/default/files/sea-astronomia_oct_2019.pdf>
- RUFUS, Z. (2023). «AI will not replace you, a person using AI will». A: *ADMARP*. 1/16/2023. <<https://admarp.com/ai-will-not-replace-you-a-person-using-ai-will>>
- RUGGERI, A. (2024). «The surprising promise and profound perils of AIs that fake empathy». A: *NewScientist*. Technology, 6/3/2024. <<https://www.newscientist.com/article/mg26134810-900-the-surprising-promise-and-profound-perils-of-ais-that-fake-empathy/>>
- SÁNCHEZ Zaplana, A. (2023). «Por qué la inteligencia artificial nos hará más humanos». A: *El Español de Alicante*. 2/3/2023. <https://www.elespanol.com/alicante/opinion/20230302/inteligencia-artificial-hara-humanos/745295475_12.html>
- SCHANK, R. (2017). «Listening to Groucho Marx make jokes can teach us a lot about AI». A: *Linkedin*. 16/6/2017. <<https://www.linkedin.com/pulse/listening-groucho-marx-makes-jokes-can-teach-us-lot-ai-roger-schank/>>
- SEMPERE, M. (2024). «Proteger el talento femenino en la era de la Inteligencia Artificial». A: *elEconomista.es*. Retail-Consumo. 7/3/2024. <<https://www.eleconomista.es/retail-consumo/noticias/12711175/03/24/proteger-el-talento-femenino-en-la-era-de-la-inteligencia-artificial.html>>
- THÉRY, M. (2012). *Isaac Asimov, Mensaje al futuro*. Documental de Keppler Productions & Arte France. <<https://www.dailymotion.com/video/x8lqiyh>>
- UNESCO (2024). «AI in education». A: *Unesco.org*. <<https://www.unesco.org/en/digital-education/artificial-intelligence>>
- W³TECHS (2024). «Usage statistics of content languages for websites». A: *World Wide Web Technology Surveys*. <https://w3techs.com/technologies/overview/content_language>

- WIKIPEDIA contributors (20 de juliol del 2024). «Symbolic artificial intelligence». A: *Wikipedia, The Free Encyclopedia*.
- WOOLDRIDGE, M. (2024). «The road to conscious machines: AI from Eliza to ChatGPT and beyond». A: *NewScientist*. 12/3/2024. <https://drive.google.com/file/d/1GuS_1vTDw9e-fhuOEnuh2HtqJfT8rkpy3/view>
- XARXA VIVES D'UNIVERSITATS (2024). *Reglament de la Lliga Debat Sènior 2024*. <https://drive.google.com/file/d/1AMws_1lU-zrkUJnIDRDXwofu0KdaQ6kTR/view>

14. IMATGES DE LA COMPETICIÓ

A continuació podem veure els protagonistes en acció, tant en unes fotos fetes pels mateixos components de l'equip Iubilo URV com en les fotos disponibles a la web de la Xarxa Vives d'Universitats, i al final tenim uns enllaços a les gravacions en vídeo dels tres debats on vam intervenir, de la mateixa web de la Xarxa Vives.

Fotos preses pels autors o pels organitzadors de la Xarxa Vives, el 5 i 6 de juny del 2024, a la competició de la Lliga Debat a la UPC, campus Nord, Barcelona.



1. Preparats per debatre, som-hi!



2. Quasi tots els competidors.



3. La Isabel i la Carme repassant.



4. L'equip de competidors de Lubilo URV a punt del primer debat.



5. La Carme: la IA sí que és una amenaça.



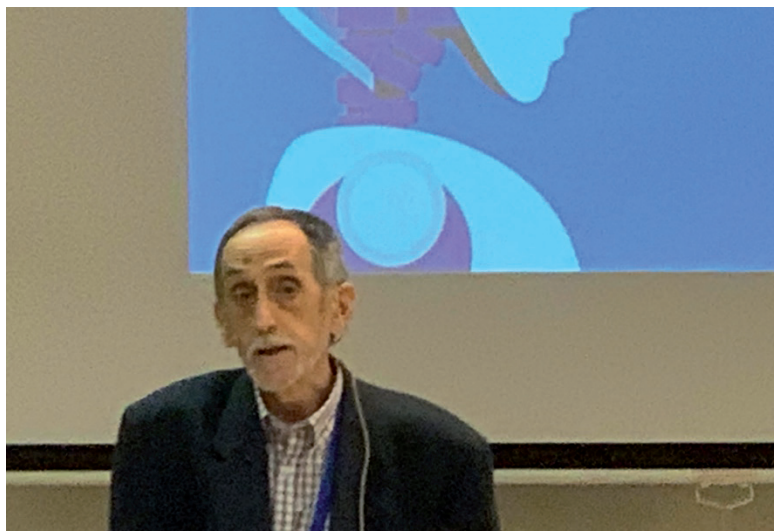
6. La Isabel: la IA sí que és una amenaça.



7. El José Eugenio: la IA sí que és una amenaça.



8. Carai, altre cop el sí!



9. L'Albert: que sí.



10. La refutadora Isabel: altre cop que sí.



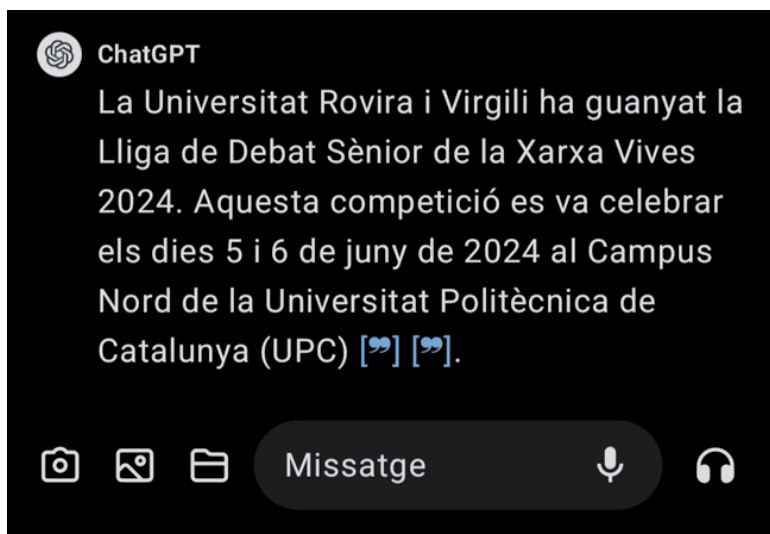
11. El Jordi: que sí.



12. Els set magnífics i el diploma de la Isabel.



13. En acabat...



14. Notícia falsa de la IA, just al final de la competició.

I aquí sota teniu els enllaços corresponents a les gravacions de vídeo dels tres debats on l'equip de Iubilo URV vam intervenir, de la web de la Xarxa Vives d'Universitats:

URV - UB 5 de juny del 2024, 10.30 h – 11.20h:

https://www.youtube.com/watch?v=XzSo4W1IEtI&list=PL05Pxx8pMlZj2yMZYaEmE85FBV2fye_Jv&index=2



URV - UPF 5 de juny del 2024, 18 h – 18.50 h:

https://www.youtube.com/watch?v=crtr3UenvR4&list=PL05Pxx8pMlZj2yMZYaEmE85FBV2fye_Jv&index=7



URV - UVic/UCC 6 de juny del 2024, 10.45 h - 11.35 h:

https://www.youtube.com/watch?v=HZn5Ua4gmCc&list=PL05Pxx8pMlZj2yMZYaEmE85FBV2fye_Jv&index=9



ELABORACIÓ DE LES FIGURES

Les imatges de les figures 4, 5, 7, 9 i 10 han estat cedides per Roger Mor Crespo i Isabel Oussedick Vilalta, i obtingudes per IA.

La imatge de la figura 3 ha estat cedida per Albert Bordons, i obtinguda per IA.

Les figures 1, 2 i 6 han estat elaborades per Albert Bordons, a partir d'imatges obtingudes per IA.

La imatge de la figura 8 és de domini públic.

OBRA DE LA COBERTA



Robert Delaunay (París 1885 - Montpellier 1941), *Rythme, Joie de vivre*, 1930. Pintura d'oli sobre tela, 200 x 228 cm. Musée National d'Art Moderne, París.

La joia o goig de viure d'aquest quadre recorda les ganes de seguir essent actius i gaudint-ne, que els autors, membres de la comunitat Iubilo URV, volen compartir amb aquesta publicació. Les figures geomètriques de la pintura també recorden les imatges generades per intel·ligència artificial, i al mateix temps suggereixen ritme, moviment i adaptabilitat.

Aquest treball recull l'experiència d'un grup multidisciplinari que va organitzar la comunitat Iubilo de la Universitat Rovira i Virgili i va competir a l'edició 2024 de la Lliga de Debat Sènior de la Xarxa Vives. L'equip va investigar, debatre i elaborar arguments al voltant de la pregunta “La intel·ligència artificial és una amenaça per a la humanitat?”, amb l'objectiu de construir una visió pròpia, rigorosa i profunda d'aquest desafiament tecnològic actual.

Cal assenyalar que els arguments presentats i debatuts ho foren a partir de les fonts d'informació disponibles durant els mesos de març a juny de 2024, i per tant alguns aspectes aquí presentats poden estar desfasats, donada la ràpida evolució d'aquest camp, del qual cada dia hi ha notícies noves.

El resultat combina la reflexió sobre els pros i contres de la intel·ligència artificial amb l'explicació de l'experiència viscuda pels autors en la preparació i la competició mateixa, amb la voluntat d'inspirar altres grups a explorar, preparar i debatre grans qüestions d'actualitat amb mirada oberta i col·laborativa.



Diputació Tarragona

