

Píndola IEI "La ciència qui la fa?". Jordi Planes entrevista a Daniel Gibert

La [secció de ciències experimentals i tecnologia](#) [

<http://www.iei-departaments.cat/ca/administracio/iei-departaments/ciencies-experimentals-matematiques-i-tecnol>] del Departament de Ciència i Tecnologia de l'Institut d'Estudis Ilerdencs (IEI) té com a prioritat divulgar el coneixement i l'experiència científica i tecnològica amb els objectius de :

- Divulgar i socialitzar la cultura científica
- Posar en valor la ciència, la tecnologia i les matemàtiques com a part del patrimoni cultural de la societat lleidatana
- Promoure les vocacions científiques i l'apropament de la ciència al jovent.

Dins aquest marc realitza una sèrie de vídeos divulgatius sota el tema: **La ciència qui la fa?** Es tracta de Píndoles que ens apropen a les persones i el seu camp de treball en l'àmbit de les ciències experimentals, les matemàtiques i la tecnologia de ponent.

En la darrera aportació de *La ciència qui la fa?* Jordi Planes Cid, professor de l'Àrea de Ciències de la Computació i Intel·ligència Artificial del Departament d'Informàtica i Enginyeria Industrial de l'Escola Politècnica Superior (EPS) de la Universitat de Lleida entrevista a Daniel Gibert Llauredó, doctor per la UdL i investigador al Centre d'Intel·ligència Artificial Aplicada a la University College Dublin.

Al 2020 en Daniel Gibert va llegir la seva Tesi Doctoral, *"Going Deep into the Cat and the Mouse Game: Deep Learning for Malware Classification"*, dirigida per Carles Mateu i el propi Jordi Planes.

Resum Tesi Doctoral:

La lluita contra el programari maliciós no s'ha interromput mai des dels inicis de l'era digital, esdevenint una carrera armamentística cíclica i interminable; a mesura que els analistes en seguretat i investigadors milloren les seves defenses, els desenvolupadors de programari maliciós continuen innovant, trobant nous vectors d'infecció i millorant les tècniques d'ofuscació. Recentment, degut al creixement massiu i continu del programari maliciós, es requereixen nous mètodes per a complementar els existents i així poder protegir satisfactòriament els sistemes de nous atacs i variants. L'objectiu d'aquesta tesi doctoral és el disseny, implementació i avaluació de mètodes d'aprenentatge automàtic per a la detecció i classificació de programari maliciós, a causa de la seva capacitat per a manipular grans volums de dades així com la seva habilitat de generalització. La recerca s'ha estructurat en quatre parts. La primera part proporciona una descripció completa dels mètodes i característiques utilitzats per a la detecció i classificació de programari maliciós. La segona part consisteix en l'automatització del procés d'extracció de característiques utilitzant tècniques d'aprenentatge profund. La tercera part consisteix en la investigació de mecanismes per a combinar múltiples modalitats o fonts d'informació per a incrementar la robustesa dels classificadors basats en aprenentatge profund. La quarta part d'aquesta tesi presenta els principals problemes i reptes als que s'enfronten els analistes en seguretat, com el problema de la desigualtat entre el nombre de mostres per família, l'aprenentatge advers, entre altres. Tanmateix, proporciona una extensa avaluació dels diferents mètodes d'aprenentatge automàtic contra diverses tècniques d'ofuscació, i analitza la utilitat d'aquestes per a augmentar el conjunt de dades d'entrenament i reduir la desigualtat de mostres per família.

