

Programma Autumn School in IA

Aula UTLC Via Vivaldi 5 Genova

Data	Titolo intervento	Docente	Abstract
10/11/2025			
Mattino 11-13	Incontro e presentazione della scuola	Elisa Bricco (Università di Genova)	
	RIC: Cyber Humanities, etica e competenze: nuovi paradigmi per la ricerca nell'era algoritmica	Giovanni Adorni (Università di Genova)	L'intervento esplora il quadro teorico e operativo delle Cyber Humanities come risposta critica e trasformativa alle sfide poste dall'intelligenza artificiale e dalle infrastrutture algoritmiche. Partendo dal Manifesto for Cyber Humanities, si analizzano i principi fondativi—dall'etica-by-design alla sovranità digitale—come base per ripensare il ruolo della ricerca umanistica in contesti computazionali. Vengono quindi discusse le competenze digitali e critiche necessarie (anche alla luce del framework europeo DigComp 3.0) per formare ricercatori e cittadini capaci di agire in modo consapevole, inclusivo e responsabile in ecosistemi ibridi uomo-macchina. Un focus particolare è dedicato all'educazione all'IA e con l'IA, con esempi di applicazione nei contesti formativi e di ricerca.
Pomeriggio 14-18	FON: Linguistica computazionale e IA generativa: un'introduzione	Rachele Sprugnoli (Università di Parma)	Il corso propone un'introduzione alla linguistica computazionale e alle applicazioni dei Large Language Model (LLM). Nella prima parte saranno illustrati i concetti e gli approcci fondamentali della disciplina, con particolare attenzione ai principali task linguistici e ai metodi utilizzati per affrontarli. Successivamente,

			l'attenzione si sposterà sugli LLM, modelli alla base di tecnologie come ChatGPT, per esplorarne i principi di funzionamento, le tecniche di prompting, le potenzialità applicative, nonché le sfide e le implicazioni sociali e culturali che sollevano.
11/11/2025			
Mattino 9-13	FON: A Theoretical and Practical Introduction to Neural Language Models: Evaluating and Exploring their Linguistic Abilities	Alessio Miaschi (ILC-CNR Pisa)	Il campo del Natural Language Processing (NLP) sta vivendo un periodo di grande progresso, dovuto in gran parte dal passaggio dagli approcci tradizionali agli algoritmi basati su reti neurali. Tra questi, i Large-scale Language Models (LLM) hanno dimostrato prestazioni eccezionali in un ampio spettro di task e nella generazione di testi coerenti e contestualmente appropriati. Tuttavia, questo progresso va di pari passo con una minore trasparenza, che rende sempre più difficile capire come tali modelli operino e quali capacità possiedano. Si offrirà una panoramica sui modelli linguistici e sui recenti progressi in questo ambito, con un focus particolare sugli studi che negli ultimi anni si sono incentrati sulle abilità linguistiche codificate implicitamente dagli LLM e su come queste intuizioni possano migliorare la nostra comprensione del comportamento di tali sistemi in diversi compiti e applicazioni. Nella seconda parte, si passerà a una sessione più pratica, fornendo ai partecipanti un'introduzione su come questi modelli possano essere utilizzati e esplorati concretamente.
Pomeriggio 14-18	Lecture: Personalizzare le intelligenze artificiali: teoria, pratica e strategie per un uso consapevole.	Alberto Puliafito (Slow News)	La personalizzazione delle intelligenze artificiali non è soltanto una questione tecnica: è un modo per renderle strumenti realmente utili al proprio lavoro di ricerca e di vita quotidiana. Partendo dalla mia esperienza diretta come giornalista, consulente e formatore, mostrerò

			come sia possibile progettare e addestrare sistemi AI personalizzati, evidenziandone i vantaggi e i limiti. La lezione sarà divisa in due parti: una panoramica teorica, che offrirà le basi per comprendere i meccanismi di personalizzazione, e una più pratica durante la quale esploreremo insieme metodi, strategie e tecniche per creare assistenti su misura. L'obiettivo è duplice: fornire strumenti concreti per integrare le AI nel proprio lavoro e stimolare una riflessione critica sul loro impatto, per imparare a usarle senza esserne usati.
12/11/2025			
Mattino 9-13	RIC. Intelligenza artificiale e applicazioni XR per il patrimonio culturale: le esperienze del laboratorio Dhekalos	Ramona Quattrini con Irene Galli - Dottoranda (Università Politecnica delle Marche)	Negli ultimi dieci anni il laboratorio Dhekalos dell'Università Politecnica delle Marche ha condotto un'intensa attività di ricerca sull'integrazione tra intelligenza artificiale e tecnologie XR per la documentazione, l'interpretazione e la valorizzazione del patrimonio culturale. La lezione propone un percorso attraverso i principali casi studio sviluppati nel laboratorio, evidenziando l'evoluzione metodologica e tecnologica delle soluzioni adottate. Si partirà dalle prime applicazioni di Deep Learning per il riconoscimento di primitive geometriche e la segmentazione semantica di nuvole di punti complesse relative all'architettura storica (Palazzo Ducale di Urbino- Progetto CIVITAS), per arrivare a progetti più recenti dedicati alla manipolazione e al morphing di immagini gigapixel finalizzati alla creazione di video-narrazioni immersive su dipinti e disegni preparatori (Disegni di Michelangelo-Dipinti di Sebastiano del Piombo- Progetto EAR; Sipario storico di Fano, Disegni di Dante Ferretti) ma anche dedicati alla guida in AR per aree archeologiche (Progetto Suasa). Saranno inoltre presentate sperimentazioni basate su modelli generativi (text-to-model e text-to-image) e Large Language Models per la costruzione di scenari immersivi fondati

			sulla conoscenza delle fonti, nonché applicazioni con chatbot e assistenti virtuali per la mediazione culturale in ambito museale (Città Ideale di Urbino). La sessione teorica sarà seguita da una demo session e da un tutorial interattivo, durante i quali i partecipanti all'autumn school potranno sperimentare direttamente alcune delle applicazioni sviluppate.
Pomeriggio 14-18	RIC: Natural Language Processing e Misoginia Digitale: Modelli, Dataset e Sfide Emergenti	Elisabetta Fersini (Università di Milano Bicocca)	La rilevazione automatica della misoginia costituisce un ambito di ricerca emergente nell'elaborazione del linguaggio naturale, con implicazioni rilevanti per l'analisi dei discorsi d'odio e per la progettazione di tecnologie linguistiche responsabili. Questo modulo intende fornire una panoramica metodologica sulla rilevazione automatica della misoginia, illustrandone i fondamenti sociali nell'era digitale, le risorse linguistiche e le principali strategie computazionali. Il modulo si articolerà in quattro parti: un'introduzione al problema della misoginia nei contesti digitali; la presentazione di benchmark di riferimento; l'analisi di approcci supervisionati e di modelli basati su machine learning, con particolare riferimento a modelli neurali; la discussione delle sfide aperte, quali la scarsità di risorse multilingue, la riduzione dei bias algoritmici e le difficoltà della misoginia multimodale.
13/11/2025			
Mattino 9-13	FON: Intelligenza Artificiale e pratiche educative: implicazioni pedagogiche e metodologiche	Mario Pireddu (Università della Tuscia)	L'integrazione dell'Intelligenza Artificiale nei contesti educativi solleva interrogativi cruciali sul piano epistemologico, didattico ed etico. Il laboratorio si propone di esplorare criticamente le trasformazioni in atto nei processi di insegnamento-apprendimento, interrogando le forme di delega cognitiva, la ridefinizione dei ruoli educativi e le logiche di automazione che caratterizzano molte delle attuali soluzioni basate sull'IA. L'intervento si articolerà in una prima parte

			teorica, volta a fornire una cornice critica sulle relazioni tra intelligenza artificiale, pedagogia e cultura digitale, seguita da una parte operativa in cui si analizzeranno scenari, strumenti e pratiche didattiche emergenti.
Pomeriggio 14-18	FON: Bias and Variation in Natural Language Processing	Alan Ramponi (FBK)	Recent developments in natural language processing (NLP), facilitated by the advent of large language models, have led to widespread adoption of language technologies across a wide range of tasks. Despite the great performance in NLP benchmarks, NLP systems developed without due care reflect and amplify harmful societal biases, leading to gender-biased outputs, greater exclusion of minorities, and misrepresentation of languages and cultures. In this lecture, I will examine the origins, manifestations, and consequences of bias in NLP. After introducing preliminary concepts for understanding how NLP and machine learning work, I will present research efforts aimed at identifying and mitigating bias in NLP and the best practices in the field, focusing specifically on fairness and inclusiveness of language and its variations. Finally, I will show how to identify biases and ensure that linguistic variation is represented in corpora using a publicly-available tool. Participants will gain a deep understanding of both the opportunities and the limitations of current NLP, and how to deal with the latter effectively.
14/11/2025			
Mattino 9-13	RIC: Il valore del disaccordo: approcci prospettivisti all'annotazione linguistica	Valerio Basile (Università di Torino)	La creazione di risorse per l'elaborazione del linguaggio naturale (NLP) coinvolge spesso partecipanti umani, "annotatori" che esprimono la loro comprensione o percezione di un dato fenomeno linguistico studiato. I giudizi forniti dagli annotatori possono avere un grande variabilità interna: il linguaggio è ambiguo e la sua comprensione soggettiva, al punto che più interpretazioni possono essere valide per lo stesso dato.

			In questa ottica, mettere a fuoco come affrontare i disaccordi nelle annotazioni è di fondamentale importanza per qualsiasi attività di raccolta dati, al fine di chiarire quando le opinioni divergenti degli esseri umani influiscono sulla qualità della risorsa finale o se rappresentano prospettive di comprensione del testo diverse (ma ugualmente valide). Questo modulo introduce l'approccio "prospettivista" alla raccolta dei dati e la modellazione di fenomeni linguistici soggettivi, una metodologia che tiene conto la diversità delle annotazioni umane. Il modulo tratta le potenziali fonti di disaccordo da considerare nella progettazione di schemi di annotazione, le strategie per valutare la qualità dei dati in un quadro prospettivistico, e i metodi per integrare il disaccordo all'interno di processi di apprendimento automatico di modelli linguistici.
13.00	Chiusura Autumn School	Cristina Bosco (Università di Torino)	