

LA SEDE LOMELLINA DEL GRUPPO DIVENTERÀ PUNTO DI RIFERIMENTO EUROPEO

Ivm Parona, leader sul mercato e nella ricerca della sostenibilità

«Mantenerci tra i gruppi più rilevanti del settore ed entrare nella cerchia dei pochi in grado di occupare la fetta principale di mercato mondiale». Gli obiettivi di IVM Parona, azienda che è già tra i leader nella produzione di vernici per legno, vengono fissati dal suo presidente, Adriano Teso.

E Parona sarà il centro di un'ulteriore evoluzione del gruppo, diventandone la sede produttiva principale in Europa e centro di ricerca di riferimento a livello mondiale. Gli attuali dipendenti passeranno dagli oltre 300 attuali a quasi il doppio.

Le vernici per legno sono un mercato particolare: migliaia i tipi di prodotto per altrettanti tipi di superficie. Una galassia che IVM cerca di soddisfare, mettendo al centro l'evoluzione tecnologica e la sostenibilità, con una ricerca mirata all'utilizzo delle materie prime, le meno impattanti possibili sull'ambiente.

«Ma la nostra attenzione - sottolinea il presidente Teso - è focalizzata prima di tutto sulla creazione di un ambiente di lavoro sicuro e salubre per i nostri dipendenti. E questo non significa semplicemente rispettare le norme legislative in vigore ma andare oltre. Questa



è la nostra prima regola». Quello della IVM è un punto di osservazione privilegiato sull'economia nazionale e mondiale. «Oggi il mercato è fatto di 8 miliardi di persone - è la visione di Adriano Teso - Ma in Italia il costo del lavoro è di 4-5 volte più alto che in altri continenti. Credo che il nostro Paese debba aprire una stagione di vere riforme strutturali, a partire da quella della giustizia, che oggi ha tempi insostenibili per le imprese sino ad arrivare alla burocrazia, che continua ad essere un grandissimo freno alla crescita». Intanto venerdì scorso alcuni studenti del corso di

LE REGOLE DI TESO

Un'immagine dall'alto della grande area sulla quale è insediata la IVM a Parona, azienda produttrice di vernici per legno. A fianco il presidente del gruppo, Adriano Teso



Brand Strategy della NABA, accompagnati dalla docente e architetto Carolina Nisivoccia, hanno trascorso una giornata presso lo stabilimento di Parona della IVM.

Dopo una breve presentazione del gruppo e delle sue caratteristiche, gli studenti hanno visitato la fabbrica e lo showrom aziendale. Successi-

vamente hanno incontrato i tecnici specializzati del laboratorio di Ricerca Ivm per un focus su finiture ed effetti speciali. «L'incontro - si legge in un comunicato - conferma l'impegno dell'azienda lombarda per una formazione di alto livello, in collaborazione con enti e poli universitari d'eccellenza.

LA CASA

ENERGETICA



ANTONIO CONSOLE
Architetto

Certificatore Energetico ed Energy Consultant

Risparmiare acqua, ecco alcuni consigli

Ci eravamo abituati a lamentarci dell'assenza delle piogge, in special modo nel Nord Italia. La siccità non è dovuta solo a una carenza di precipitazioni molto estesa nel tempo, ma anche a temperature più alte della norma che, in associazione al prolungato bel tempo, hanno portato a un aumento della quantità d'acqua che evapora dal terreno e traspira dalle piante (evapotraspirazione).

Così è stato sino all'altro ieri, in special modo nell'Emilia Romagna. Stiamo imparando che, a causa dei cambiamenti climatici abbiamo lunghi periodi di bel tempo e brevi intensi tempi di piogge. La quantità d'acqua che proviene dal cielo non è cambiata: è cambiato il modo. Lunghi periodi di siccità spezzati da intensi acquazzoni che devastano il territorio senza lasciar benefici nelle nostre campagne. Per affrontare questi problemi, vi sono programmi di conservazione dell'acqua, miglioramenti delle infrastrutture e un maggiore monitoraggio e gestione delle risorse idriche. Dobbiamo ricordarci che l'Italia è un giardino: è nostro compito farlo rimanere un giardino. La riduzione delle emissioni di gas serra, insieme al passare a fonti di energia rinnovabili, può mitigare gli impatti dei cambiamenti climatici. Ma cosa possiamo fare per mitigare questi problemi, oltre a consumare meno energia? Ci sono diversi modi per risparmiare acqua, in special modo quella più preziosa: l'acqua potabile. Ecco alcuni suggerimenti. Installare rubinetti e soffioni a basso flusso. Riparare tempestivamente le perdite. Installare un WC a doppio scarico. Usare la lavastoviglie: consuma meno acqua rispetto al lavaggio manuale dei piatti. Innaffiare le piante con saggezza: presto al mattino o alla sera. Installare un sistema di raccolta dell'acqua piovana. E aggiungo: utilizzare l'acqua fredda e meno possibile l'acqua calda. Potete scrivermi a tony.console07@gmail.com.

COMPUTER TECHNOLOGY PLANET



Come funzionano i modelli linguistici di ChatGPT?

di Paolo Vella - paolovella1993@gmail.com

Chiunque padroneggi a sufficienza la sua lingua madre è in grado di indovinare, con

alte probabilità di successo, quale sia la parola che conclude una determinata frase o periodo. Un large language model (Llm), un modello linguistico ampio (di quelli alla base di strumenti di intelligenza artificiale come ChatGPT) in sintesi estrema, non fa che replicare questa operazione per via informatica, potendola potenzialmente proseguire anche all'infinito. Per esempio, quando fornite a ChatGPT una frase come input e poi gli chiedete di proseguire nella scrittura, il sistema di OpenAI non farà altro che estrarre dal suo database le parole o le frasi che hanno la maggior probabilità di essere coerenti con quelle che le hanno precedute, arrivando anche a scrivere interi articoli o addirittura saggi (anche se la qualità, per il momento, lascia spesso a desiderare). Più nello specifico, come si legge su AI Multiple, "un large language model è un modello basato su machine learning addestrato su un vasto corpus di testi, allo scopo di generare output in vari ambiti della elaborazione naturale del linguaggio (natural language processing, Nlp) come la generazione di testi, rispondere alle domande e la traduzione automatica". Prima di vedere meglio come questi sistemi vengono addestrati, è il caso di spiegare perché sono considerati "large". A determinare la loro dimensione è il numero di "parametri" o "pesi" impiegati all'interno di queste reti neurali. I parametri, a loro volta, rappresentano le connessioni che si creano tra i vari "nodi" presenti nelle reti neurali. Per fare un parallelismo con il cervello umano (sul quale le reti neurali sono approssimativamente modellate), si potrebbe dire che i nodi sono l'equivalente dei nostri neuroni, mentre i



parametri sono l'equivalente delle nostre sinapsi. Di conseguenza, la quantità dei parametri di un modello fornisce una buona approssimazione di quanto sarà accurato il lavoro da esso svolto (anche se ci sono altri fattori in ballo). Questi modelli linguistici sono quindi "large" perché al loro interno sono presenti miliardi, e nei sistemi più recenti anche centinaia o migliaia di miliardi, di parametri. Per esempio, il noto GPT-3 (il sistema che alimenta ChatGPT e che adesso sta venendo sostituito con l'ancor più grande GPT-4) ha 175 miliardi di parametri, mentre MT-NLG di Nvidia e Microsoft arriva a 530 miliardi. Il più grande in assoluto è però WuDao 2.0 dell'Accademia di Pechino per l'intelligenza artificiale, dotato della strabiliante quantità di 1.750 miliardi di parametri. I large language model sono quindi reti neurali estremamente vaste impiegate a scopi linguistici, e di conseguenza rappresentano una forma di algoritmo di deep learning in grado di riconoscere, riassumere, tradurre,

prevedere e generare testi e altri contenuti sulla base della conoscenza appresa dal loro dataset.

Non è tutto, perché - a differenza dei più diffusi algoritmi di deep learning del recente passato, basati su tecniche note come convolutional neural network o recursive neural network - gli Llm sono un sottoinsieme dei cosiddetti transformer model, introdotti per la prima volta da Google nel 2017. Un transformer (come sono più comunemente noti) è una rete neurale che apprende il modo in cui i dati vengono utilizzati tenendo traccia delle relazioni all'interno delle sequenze che li contengono. È il caso delle parole contenute in una frase, ma anche del linguaggio informatico usato per scrivere codice o della sequenza delle proteine (com'è il caso del sistema AlphaFold, sempre di Google, in grado di prevedere la struttura delle proteine).

"Prima dell'arrivo dei transformer, i programmatori dovevano addestrare le reti neurali tramite enormi database etichettati (segnalando all'algoritmo, per esempio, che cos'è presente nelle immagini, ndr), che sono costosi da produrre e richiedono moltissimo tempo - si legge sul sito di Nvidia -. Poiché trovano le correlazioni tra gli elementi per via matematica, i transformer eliminano questo bisogno, rendendo utilizzabili le migliaia di miliardi di immagini e i petabytes di dati testuali presenti sul web".

I dati utilizzati per l'addestramento degli Llm non sono infatti stati precedentemente etichettati, di conseguenza la fase preparatoria all'addestramento è molto più agevole e soprattutto è possibile utilizzare colossali quantità di dati (nel caso di GPT-3, per esempio, sono stati impiegati 800 gigabytes di informazioni, tra cui l'intera Wikipedia in lingua inglese), che sarebbe impossibile etichettare a mano.