

این وبینار در تاریخ ۲۸ مرداد ماه برگزار گردید. مدل های زبانی بزرگ (LLM ها) به دلیل ماهیت داده محور و پیچیدگی ساختاری، در معرض حملات خصمانه و ریسک های افشای اطلاعات می باشند. بر این اساس، در ابتدای کارگاه به انواع حملات خصمانه و ریسک های افشای اطلاعات در LLM ها پرداخته شد. در ادامه خطرات استخراج شده توسط HIPAA برای مدل های زبانی بزرگ در حوزه سلامت و نشت اطلاعات حساس بیماران مورد اشاره قرار گرفت و در انتها ، توصیه های HIPAA با عنوان بهترین شیوه های امنیتی برای کاهش این مخاطرات ارائه شد.