



تحلیل داده‌های ژئوشیمیایی با یادگیری ماشین بدون نظارت

۳	<u>پس‌زمینه و ضرورت پروژه</u>
۴	<u>چالش‌های فعلی و نیاز به نوآوری</u>
۴	<u>هدف و چشم‌انداز پروژه</u>
۴	<u>رویکرد و مسیر حل مسئله</u>
۵	<u>ساختار و جریان پروژه</u>
۵	<u>ارزش و تأثیر پروژه</u>
۷	<u>معماری راه‌حل و رویکرد فناورانه</u>
۷	<u>قیود اجرایی و چالش‌های عملی</u>
۸	<u>مراحل اجرایی پروژه</u>
۸	<u>شاخص‌های موفقیت و معیارهای ارزیابی</u>
۹	<u>چشم‌انداز و توسعه آینده</u>

پس‌زمینه و ضرورت پروژه

اکتشاف معدن، ترکیبی از علم، تجربه و داده است؛ اما در دهه اخیر، داده نقش پررنگ‌تری از هر زمان دیگری یافته است. در میان داده‌های اکتشافی، داده‌های ژئوشیمی جایگاه ویژه‌ای دارند، زیرا ترکیب شیمیایی نمونه‌های سطحی و زیرسطحی می‌تواند سرنخ‌هایی از وجود کانی‌زایی در عمق زمین ارائه دهد. با این حال، تفسیر داده‌های ژئوشیمیایی همواره یکی از پیچیده‌ترین مراحل اکتشاف بوده است. حجم بالای متغیرها، اثر لیتولوژی، تفاوت مقیاس‌ها (wt, ppm) و حضور نویز، کار را برای تحلیل سنتی بسیار دشوار می‌کند.

در پروژه‌های معدنی امروزی، تصمیم‌گیری دیگر تنها بر مبنای شهود و تجربه ممکن نیست. هر خطای تحلیلی می‌تواند میلیون‌ها تومان هزینه و ماه‌ها زمان به پروژه تحمیل کند. در چنین شرایطی، نیاز به روش‌هایی نو احساس می‌شود؛ روش‌هایی که بتوانند از دل داده‌های خام و بدون تفسیر انسانی، الگوهای پنهان را آشکار کنند. این پروژه در پاسخ به همین نیاز شکل گرفت؛ تلاشی برای بهره‌گیری از یادگیری ماشین بدون نظارت، برای استخراج معنا از داده‌های ژئوشیمیایی خام.

چالش‌های فعلی و نیاز به نوآوری

روش‌های سنتی تحلیل ژئوشیمی، مانند نمودارهای دو یا سه‌عنصری، تحلیل خوشه‌ای بصری، یا تکیه بر آستانه‌های تجربی، اگرچه دهه‌ها کاربرد داشته‌اند، اما در مواجهه با حجم زیاد داده و روابط پیچیده چندعنصری کارایی محدودی دارند. این روش‌ها معمولاً به تفسیر ذهنی وابسته‌اند و نتایج آن‌ها از فردی به فرد دیگر متفاوت است.

داده‌های ژئوشیمی، برخلاف تصور، همیشه تمیز نیستند. وجود مقادیر پرت (outlier)، مقادیر گم‌شده (مانند ۹۹۹-) و تفاوت چند مرتبه‌ای در مقیاس عناصر مختلف، باعث می‌شود مدل‌های آماری کلاسیک توانایی توصیف ساختار داده را نداشته باشند. از سوی دیگر، در داده‌های اکتشافی معمولاً نسبت نقاط آنومالی به نقاط زمینه بسیار پایین است؛ یعنی با پدیده‌ای ذاتاً نامتوازن (imbalanced) روبه‌رو هستیم.

در چنین شرایطی، استفاده از مدل‌های یادگیری نظارتی (Supervised Learning) عملاً ممکن نیست، زیرا برچسب واقعی برای "آنومالی" وجود ندارد. بنابراین، تنها گزینه علمی و منطقی، استفاده از روش‌های یادگیری بدون نظارت است. رویکردی که به مدل اجازه می‌دهد خود، الگوهای طبیعی داده را کشف کرده و نقاط متفاوت را به‌عنوان آنومالی معرفی کند.

هدف و چشم‌انداز پروژه

هدف اصلی پروژه، توسعه یک چارچوب تحلیلی مبتنی بر یادگیری ماشین بود که بتواند بدون نیاز به برچسب‌های انسانی، رفتار طبیعی داده‌های ژئوشیمیایی را یاد بگیرد و از طریق مقایسه الگوها، مناطق غیرعادی (آنومالی) را شناسایی کند.

این پروژه با هدف افزایش دقت، کاهش ریسک و تسریع تصمیم‌گیری‌های اکتشافی طراحی شد. در واقع، ما می‌خواستیم فراتر از "تحلیل داده" حرکت کنیم و به مرحله‌ای برسیم که داده بتواند خود درباره خودش سخن بگوید.

چشم‌انداز نهایی پروژه، ساخت یک مدل پایدار و تعمیم‌پذیر بود که بتواند در پروژه‌های مختلف و با داده‌های متفاوت نیز به‌کار رود. چنین مدلی می‌تواند به ابزار تصمیم‌یار زمین‌شناسان و مدیران اکتشاف تبدیل شود؛ ابزاری که به جای جایگزینی نیروی انسانی، بینش آن‌ها را تقویت می‌کند.

رویکرد و مسیر حل مسئله

ما این پروژه را با یک اصل ساده آغاز کردیم: داده باید تمیز، مقیاس پذیر و قابل تفسیر باشد. نخستین گام، پاک سازی کامل داده ها بود. مقادیر غیرواقعی (مانند ۹۹۹-) به مقادیر گمشده تبدیل شدند، ستون های کاملاً تهی حذف شدند و توزیع آماری هر عنصر بررسی شد. پس از آن، داده ها نرمال و باز مقیاس شدند تا عناصر با واحدهای متفاوت در یک چارچوب تحلیلی مشترک قرار گیرند.

در مرحله دوم، داده ها به دو گروه تقسیم شدند: عناصر اصلی (Major oxides) و عناصر ردیابی (Trace elements). برای عناصر اصلی از تبدیل ترکیبی (CLR) و برای عناصر ردیابی از مقیاس گذاری مقاوم (Robust Scaler) استفاده شد. این تمایز از نظر زمین شیمیایی حیاتی بود، زیرا از سلطه اکسیدهای درصدی مانند SiO_2 یا Al_2O_3 بر رفتار آماری سایر عناصر جلوگیری می کرد.

در مرحله سوم، چندین الگوریتم unsupervised مورد آزمایش قرار گرفت:

Isolation Forest برای شناسایی نقاط پرت بر اساس درخت تصمیم، Local Outlier Factor برای بررسی چگالی همسایگی، و Autoencoder برای یادگیری ساختار چندبعدی داده ها و تشخیص آنومالی بر اساس خطای بازسازی.

در نهایت، Autoencoder بهترین عملکرد را نشان داد. این مدل با آموزش فقط بر داده های پس زمینه، توانست تفاوت میان داده های زمینه و آنومالی را با نسبت خطای بازسازی ۷۵/۲ برابری تشخیص دهد. نشانه ای قوی از درک ساختار واقعی داده ها توسط مدل.

ساختار و جریان پروژه

پروژه به صورت گام به گام طراحی شد تا در هر مرحله، شناخت عمیق تری از داده حاصل شود.

در فاز نخست، داده ها جمع آوری، پاک سازی و نرمال شدند. در فاز دوم، مدل های پایه اجرا شدند تا رفتار اولیه داده ها سنجیده شود. فاز سوم شامل مهندسی ویژگی و تنظیم مدل Autoencoder بود، که در آن از ساختار باریک و denoising برای جلوگیری از overfitting استفاده شد. در فاز چهارم، تحلیل نتایج و تفسیر زمین شیمیایی انجام گرفت.

هر فاز، علاوه بر خروجی فنی، بینش جدیدی نیز برای فاز بعدی فراهم کرد. این چرخه تکرارشونده باعث شد مدل به صورت پویا با داده‌ها سازگار شود و درک واقعی‌تری از ساختار آن‌ها پیدا کند.

ارزش و تأثیر پروژه

ارزش این پروژه در دو سطح علمی و اقتصادی قابل بیان است. از منظر علمی، این پروژه نشان داد که داده‌های ژئوشیمی خام را می‌توان با ابزارهای یادگیری ماشین، بدون نیاز به برچسب، به بینش‌های قابل تفسیر تبدیل کرد. این دستاورد، گامی مهم در جهت استفاده از علم داده در اکتشاف مدرن است.

از منظر اقتصادی، مدل توسعه‌یافته می‌تواند نقش قابل‌توجهی در کاهش هزینه‌های حفاری ایفا کند. در شرایطی که هر حفاری اشتباه هزینه‌ای سنگین دارد، توانایی مدل در اولویت‌بندی هوشمند نقاط با پتانسیل بالا، به معنای صرفه‌جویی مستقیم در سرمایه و زمان است. همچنین، کاهش زمان تحلیل از چند هفته به چند روز، امکان تصمیم‌گیری سریع‌تر را فراهم می‌کند. مزیتی کلیدی در پروژه‌های معدنی رقابتی.

فراتر از این، مدل ما با افزایش شفافیت و قابلیت تکرار نتایج، اعتماد تیم‌های فنی و مدیریتی را تقویت می‌کند؛ چرا که تصمیمات بر پایه تحلیل داده‌محور و قابل دفاع گرفته می‌شوند.

معماری راه حل و رویکرد فناورانه

معماری پروژه بر اساس سه اصل طراحی شد: سادگی، تفسیرپذیری و پایداری. داده‌ها پس از پیش‌پردازش وارد یک ساختار Autoencoder سه‌لایه شدند. مدل تنها با داده‌های پس‌زمینه آموزش دید و هدف آن بازسازی دقیق داده‌های "عادی" بود. هر نمونه‌ای که بازسازی آن با خطای زیاد همراه می‌شد، به عنوان آنومالی شناسایی شد.

برای جلوگیری از یادگیری بیش‌ازحد، مدل به صورت denoising طراحی شد، یعنی با داده‌های دارای نویز آموزش دید تا بتواند ساختار اصلی داده‌ها را تشخیص دهد. استفاده از منظم‌سازی (Regularization) و انتخاب گلوگاه باریک در شبکه، باعث شد مدل فقط الگوهای واقعی زمینه را یاد بگیرد، نه جزئیات تصادفی.

تمام مراحل در محیط Python و با کتابخانه‌های استاندارد (matplotlib، scikit-learn، pandas) انجام شد تا قابلیت بازتولید و انتقال به پروژه‌های دیگر حفظ شود.

قیود اجرایی و چالش‌های عملی

یکی از مهم‌ترین محدودیت‌های پروژه، نبود داده برچسب‌دار بود. بدون دانستن اینکه کدام نمونه‌ها واقعاً آنومالی هستند، ارزیابی عملکرد مدل دشوار است. به همین دلیل، تحلیل کیفی و تفسیر زمین‌شیمیایی جایگزین معیارهای کلاسیک ارزیابی شد.

چالش دیگر، کیفیت متغیر داده‌ها بود. بخشی از داده‌ها دارای نویز یا خطاهای اندازه‌گیری بودند که نیاز به پاک‌سازی دقیق داشت. همچنین، اختلاف زیاد در واحدها و دامنه عناصر باعث می‌شد مدل در مراحل اولیه به سمت عناصری با مقدار زیاد (مانند Si یا Ca) تمایل پیدا کند، که با نرمال‌سازی ترکیبی برطرف شد.

در نهایت، تفسیرپذیری نتایج نیز چالشی کلیدی بود. مدل هرچند بدون نظارت آموزش دید، اما لازم بود نتایج آن از نظر زمین‌شیمیایی معنا داشته باشند. این هدف با تحلیل فیچری و مقایسه میانه عناصر در گروه آنومالی و زمینه محقق شد.

مراحل اجرایی پروژه

پروژه در پنج فاز اجرایی پیاده‌سازی شد:

فاز اول: گردآوری، پاک‌سازی و آماده‌سازی داده‌های ژئوشیمی (تبدیل مقادیر پرت، حذف ستون‌های تهی و نرمال‌سازی).

فاز دوم: طراحی مدل‌های اولیه unsupervised و ارزیابی آماری اولیه.

فاز سوم: بهینه‌سازی Autoencoder با ساختار باریک و denoising.

فاز چهارم: تحلیل فیچری و تفسیر زمین‌شیمیایی نتایج.

فاز پنجم: تدوین شاخص‌های تصمیم‌یار و طراحی پیشنهادات توسعه آینده.

هر فاز با مستندسازی کامل همراه بود تا در پروژه‌های مشابه به‌عنوان مرجع استفاده شود.

شاخص‌های موفقیت و معیارهای ارزیابی

در نبود داده‌های برجسب‌دار، موفقیت مدل با شاخص‌های غیرمستقیم سنجیده شد.

نخست، نسبت خطای بازسازی کل به پس‌زمینه (۷۵/۲) به‌عنوان شاخص تمایز در نظر گرفته شد. عددی که نشان می‌دهد مدل تفاوت قابل‌توجهی بین زمینه و آنومالی قائل شده است.

دوم، پایداری مدل در آموزش‌های مکرر بررسی شد؛ مقدار ۹۸/۰ نشان داد مدل تقریباً همواره به نتایج مشابه می‌رسد.

سوم، تطابق زمین‌شیمیایی نتایج، یعنی هم‌خوانی عناصر آنومالی با ترکیبات شناخته‌شده کانی‌زایی منطقه. ترکیب عناصر Pb-Zn-Sb-S در میان ۱۵ عنصر برتر، تأییدکننده درستی رفتار مدل بود.

در کنار این شاخص‌ها، زمان تحلیل و میزان کاهش داده‌های غیرمفید نیز به‌عنوان معیارهای موفقیت در نظر گرفته شد.

چشم‌انداز و توسعه آینده

این پروژه تنها آغاز راه است. با افزایش حجم داده‌های ژئوشیمی و ترکیب آن با داده‌های ژئوفیزیکی، می‌توان مدل‌های چندمنبعی (Multimodal) توسعه داد که دید جامع‌تری از سیستم‌های کانی‌زایی ارائه دهند.

در فازهای بعدی، افزودن داده‌های حفاری و نتایج آزمایشگاهی می‌تواند امکان استفاده از مدل‌های نیمه‌نظارتی (Semi-Supervised Learning) را فراهم کند. همچنین، طراحی داشبوردهای تحلیلی برای نمایش آنومالی‌ها، اولویت‌بندی اهداف حفاری و گزارش‌های مدیریتی از جمله گام‌های آینده است.

فراتر از هر فناوری، ارزش واقعی این پروژه در تغییری است که در نگرش اکتشافی ایجاد می‌کند:

اکتشاف دیگر صرفاً جست‌وجوی ماده معدنی نیست، بلکه کشف دانش از داده‌هاست.