

برای تشخیص کارت‌های علامت‌گذاری‌شده، این امر مستلزم مداخله جزئی اپراتورهای انسانی حاضر در اتاق بوده بدین ترتیب که حین تلاش دستگاه برای شناسایی حروف، یک تکنیسین صحت یا عدم‌صحت تشخیص آن را اعلام می‌کرد، اما درنهایت مارک ۱ از نتایج صحیح و اشتباهات خود الگومی گرفت و الگوهای مشخص‌کننده ویژگی‌های بصری حروف، نظیر خط‌مورب حرف A یا انحنای مضاعف حرف B را با دقت تعیین می‌کرد، روزنبلات هنگام نمایش این دستگاه، روشی برای اثبات ماهیت اکتسابی این رفتار داشت، او با دستکاری اتصالات داخلی قفسه‌های تجهیزات الکتریکی و قطع اتصال برخی سیم‌ها که نقش **نورون** های مصنوعی را ایفا می‌کردند، عملکرد دستگاه را مختل می‌کرد، پس از برقراری مجدد اتصالات، دستگاه مجدداً در تشخیص حروف با مشکل مواجه‌می‌شد، اما پس از پردازش نمونه‌های بیشتر و یادگیری مجدد همان مهارت، عملکرد قبلی خود را بازمی‌یافت. عملکرد نسبتاً مناسب این ابزار الکتریکی موجب جلب‌توجه محافل علمی و صنعتی فراتر از نیروی دریایی شد. در سال‌های آتی، مؤسسه تحقیقات استنفورد (اس‌آرای)، یک مرکز پژوهشی در شمال کالیفرنیا، به مطالعه و توسعه همین مفاهیم پرداخت و آزمایشگاه شخصی روزنبلات نیز موفق به انعقاد قراردادهایی با سرویس پستی ایالات‌متحده ونیروی هوایی این کشور شد. سرویس پستی نیازمند راهکاری جهت خواندن نشانی‌های پستی روی پاکت‌ها بود و نیروی هوایی نیز درصدد شناسایی اهداف در تصاویر هوایی برآمد. با این حال، تمام این کاربردها در آینده‌ای نزدیک متصور بودند. سیستم ابداعی روزنبلات در خواندن حروف چاپی که وظیفه‌ای نسبتاً سبیل به‌شمار می‌رفت، تنها از کارایی محدودی برخوردار بود. در فرآیند تحلیل کارت‌های حاوی حرف A، هر فوتوسل ناحیه‌معینی از کارت را مورد ارزیابی قرار می‌داد؛ به‌عنوان مثال، ناحیه‌ای در مجاورت گوشه پایین سمت راست، درصورتی‌که میزان تیرگی در آن ناحیه نسبت به روشنی غالب بود، مارک ۱ به آن ناحیه «وزن» بالایی اختصاص می‌داد، بدین معنا که آن ناحیه نقش مؤثرتری در محاسبات ریاضی نهایی که تعیین‌کننده هویت حرف A بودند، ایفا می‌کرد. هنگام پردازش یک کارت جدید، دستگاه در صورتی قادر به تشخیص حرف A بود که اغلب نواحی با وزن بالا به رنگ سیاه ظاهر می‌شدند؛ این محدوده توانایی این فناوری در آن زمان بود. این فناوری از جابجایی لازم جهت تشخیص بی‌نظمی‌های موجود در ارقام دست‌نویس بهره‌مندنبود.

با وجود محدودیت‌های مشهود این سیستم، روزنبلات همچنان دیدگاهی خوشبینانه‌نسبت به‌آینه‌ان داشت. سایرین نیز معتقدبودند که این فناوری در سال‌های پیش رو ارتقاء یافته‌و قادر به فراگیری وظایف پیچیده‌تر با رویکردهای تکامل یافته‌تری خواهد بود. با این حال، این مسیر بایک چالش اساسی مواجه‌شد؛ «ماروین مینسکی».

چالش‌های یک نام‌گذاری

فرانک روزنبلات و ماروین مینسکی از هم‌دوره‌ای‌های دبیرستان علوم برانکس محسوب می‌شدند. در سال ۱۹۴۵، والدین مینسکی او را به «فیلیپس آندور»، یک مدرسه‌مقدماتی نمونه در ایالات‌متحده منتقل کردند. پس از خاتمه جنگ جهانی دوم، او در دانشگاه هاروارد به تحصیل پرداخت. با این وجود، او معتقد بود که هیچ کدام از این مراکز آموزشی با تجربه‌اش در دبیرستان علوم برانکس قابل مقایسه نبودند. جایی که به‌زعم او، دروس از چالش بیشتری برخوردار بوده و دانش آموزان بلندپروازتر بودند.

او دراین باره گفت: «افرادی که‌می‌توانستی پیچیده‌ترین ایده‌های خود را با آنها به‌بحث بگذاری و هیچ‌کس رفتاری از سر تحقیر نشان نمی‌داد.» پس از درگذشت روزنبلات، مینسکی از هم‌کلاسی سابق خود به‌عنوان نمونه‌ای از متفکران خلاق‌ی یاد کرد که در راه‌روهای دبیرستان علوم برانکس تردد‌داشتند. مینسکی نیز، مشابه روزنبلات، از پیشگامان حوزه هوش مصنوعی به‌شمار می‌رفت، اما رویکرد او نسبت به این زمینه متفاوت بود.

در دوره کارشناسی در دانشگاه هاروارد، مینسکی با بهره‌گیری از بیش از سه هزار لامپ خلأ و قطعاتی از یک بمب‌افکن قدیمی بی-۲، دستگاهی ساخت که احتمالاً نخستین شبکه عصبی مصنوعی محسوب می‌شود و آن را «Stochastic Neural Analog» (SNARC Reinforcement Calculator) نامید. سپس در اوایل دهه ۱۹۵۰ میلادی و در دوره تحصیلات تکمیلی، او به‌پیگیری مفاهیم ریاضیاتی پرداخت که نهایتاً به ظهور پرسپترون منجر شد. با وجود این، او هوش مصنوعی را یک عرصه علمی وسیع‌تر می‌دانست. او در زمره معدود دانشمندانی بود که در گردهمایی کالج دارتموث در تابستان سال ۱۹۵۶، هوش مصنوعی را به‌عنوان یک حوزه مطالعاتی مستقل و نظام‌مند تثبیت کردند. پروفسوری از دارتموث به‌نام «جان مک‌کارتی»، از جامعه علمی خواست تا به بررسی حوزه پژوهشی‌ای بپردازند که او آن را «مطالعات انوماتا» می‌نامید، اما این اصطلاح برای بسیاری از افراد فاقد معنای واضح بود. از این‌رو، او عنوان آن را به «هوش مصنوعی» تغییر داد و در همان تابستان کنفرانسی را با همکاران چندنفراز دانشگاهیان و پژوهشگران هم‌فکر سازمان دهی کرد.

دستور کار کنفرانس تابستانی تحقیقات هوش مصنوعی دارتموث، علاوه بر «شبکه‌های عصبی» موضوعاتی نظیر «ریانه‌های خودکار»، «انتزاع‌ا و «خودپه‌سازی» را نیز در بر می‌گرفت. شرکت‌کنندگان در این کنفرانس، رهبری این جریان علمی را در دهه ۱۹۶۰ میلادی عهده‌دار شدند که از برجسته‌ترین آنها می‌توان به جان مک‌کارتی که نهایتاً تحقیقات خود را به دانشگاه استنفورد در ساحل غربی منتقل کرد؛ «هربرت سایمون» و «آلن نیول»، بنیان‌گذاران آزمایشگاهی در دانشگاه کارنگی ملون واقع در پیتسبورگ و ماروین مینسکی که در مؤسسه فناوری ماساچوست در نیوا انگلند مستقر شد، اشاره کرد.

هدف این دانشمندان، بازسازی هوش انسانی با استفاده از تمام فناوری‌های موجود بود و آنها مطمئن بودند که این کار به‌زودی انجام خواهد شد؛ حتی برخی پیش‌بینی می‌کردند که یک ماشین تا ۱۰ سال دیگر می‌تواند قهرمان شطرنج جهان را شکست دهد و یک

قضیه ریاضی جدید کشف کند. مینسکی که از همان دوران جوانی موهایش ریخته‌بود، گوش‌های پهنی داشت و لبخند شیطنت‌آمیزش جلب‌توجه می‌کرد، به یک مُبلغ پرشور هوش مصنوعی بدل شد، اما این تشبیر شامل شبکه‌های عصبی نمی‌شد. از نظر او و بسیاری از همکارانش، شبکه عصبی صرفاً یکی از رویکردهای ساخت هوش مصنوعی به‌شمار می‌رفت و آنها به‌تدریج به کاوش در مسیرهای تحقیقاتی دیگر پرداختند. در اواسط دهه ۱۹۶۰ میلادی، با معطوف‌شدن توجه مینسکی به سایر روش‌های هوش مصنوعی، او در این مورد تردید کرد که آیا شبکه‌های عصبی می‌توانند از پس وظایفی فراتر از آنچه روزنبلات در آزمایشگاه خود در شمال ایالت نیویورک به‌نمایش گذاشته بود، برآیند یا خیر.

مینسکی در زمره منتقدان گسترده‌تر ایده‌های روزنبلات قرار داشت. چنان‌که خود روزنبلات در کتاب سال ۱۹۶۲ خود با عنوان «اصول نورودینامیک» اشاره کرد، پرسپترون در میان محافل دانشگاهی به‌عنوان یک مفهوم مناقشه‌برانگیز تلقی می‌شد و او بخش عمده‌ای از این وضعیت را ناشی از عملکرد رسانه‌های دانست. به باور روزنبلات، خبرنگاران در اواخر دهه ۱۹۵۰ با شور و هیجانی بی‌پروا، مانند دسته‌ای از سگ‌های شکاری خوشحال، به پوشش دستاوردهای او پرداختند. او دوباره از زبته‌هایی نظیر آنچه در اکلاهما منتشر شده بود، ابراز ناخرسندی و تصریح کرد که چنین رویکردهایی تا ایجاد اعتماد به کار او به‌عنوان یک تلاش علمی جدی، فاصله‌بسیار زیادی داشتند.

چهار سال پس از واقعه واشنگتن، روزنبلات با تعدیل ادعاهای پیشین خود تأکید کرد که پرسپترون تلاشی در راستای توسعه هوش مصنوعی به مفهوم مورد نظر محققانی نظیر مینسکی نبوده است. او در این باره می‌نویسد: «برنامه پرسپترون اساساً معطوف به ابداع دستگاه‌هایی برای «هوش مصنوعی» نیست، بلکه هدف اصلی آن بررسی ساختارهای فیزیکی و اصول نورودینامیکی است که زیربنای «هوش طبیعی» را تشکیل می‌دهند.» او در ادامه می‌افزاید: «کاربرد آن در توانمندسازی ما برای تعیین شرایط فیزیکی لازم جهت ظهور ویژگی‌های متنوع روان‌شناختی نهفته است.» به بیان دیگر، هدف او درک مکانیسم عملکرد مغز انسان بود، نه خلق یک مغز مصنوعی و عرضه آن به جهان. از آنجا که مغز ساختاری پیچیده و ناشناخته بود، امکان بازآفرینی دقیق آن وجود نداشت. با این حال، او بر این باور بود که می‌توان از ماشین‌ها به‌منظور کاوش در این معمای پیچیده و احتمالاً یافتن راه‌حلی برای آن بهره جست.

از همان آغاز، تمایزات بین هوش مصنوعی و حوزه‌هایی نظیر علوم کامپیوتر، روان‌شناسی و علوم اعصاب مبهم بود، زیرا با ظهور این فناوری نوظهور، گرایش‌های آکادمیک متعددی شکل گرفت که هر یک، چشم‌انداز این حوزه را به شیوه خاص خود ترسیم می‌کردند. برخی روان‌شناسان، متخصصان علوم اعصاب و حتی دانشمندان علوم کامپیوتر، ماشین‌ها را با دیدگاهی مشابه روزنبلات، به‌عنوان انعکاسی از عملکرد مغز تلقی می‌کردند. در مقابل، گروهی دیگر با نگرشی انتقادی به این ایده بلندپروازانه می‌نگریستند و استدلال می‌کردند که سازوکار عملکرد رایانه‌ها هیچ شباهتی به مغز ندارد و در صورت تلاش برای تقلید هوش، باید این امر را از طریق مکانیسم‌های منحصر به‌فرد خود به انجام رساند. با این حال، در آن زمان، هیچ‌یک از محققان به ساخت سیستمی که بتوان به‌درستی آن را «هوش مصنوعی» نامید، نزدیک نشده بودند. با وجود تصور بنیان‌گذاران این حوزه مبنی بر کوتاهی مسیر بازآفرینی مغز، واقعیت حاکی از طولانی‌بودن این راه بود. اشتباه اساسی آنها در نام‌گذاری حوزه پژوهشی خود به‌عنوان «هوش مصنوعی» نهفته بود. این نام‌گذاری برای دهدها این تصور را در ناظران ایجاد کرد که دانشمندان در آستانه بازآفرینی توانایی‌های مغز قرار دارند، در حالی‌که در عمل چنین نبود.

جنجال مینسکی

در سال ۱۹۶۶ میلادی، حدود چنده پژوهشگر به پورتوریکو سفر کرده و در هتل هیلتون واقع در سن خوان گرد هم آمدند. هدف از این گردهمایی، تبادل نظر پیرامون آخرین دستاوردهای حوزه موسوم به «تشخیص الگو» بود. فناوری‌ای که قابلیت شناسایی الگوها در تصاویر و سایر داده‌ها را دارا بود. در حالی‌که روزنبلات، پرسپترون را به‌عنوان مدلی از عملکرد مغز تلقی می‌کرد، سایرین آن را ابزاری در خدمت تشخیص الگومی دانستند. در سال‌های آتی، برخی مفسران تصور می‌کردند که روزنبلات و مینسکی در کنفرانس‌های علمی نظیر کنفرانس سن خوان، به مناظره‌های علنی پیرامون آینده پرسپترون می‌پرداختند، اما رقابت میان آنها بیشتر جنبه ضمنی داشت. روزنبلات اساساً به پورتوریکو سفر نکرد. در داخل هتل هیلتون، زمانی تنش آشکار شد که دانشمند جوانی به‌نام «جان مونسان» در کنفرانس به ایراد سخنرانی پرداخت. مونسان در مؤسسه تحقیقات استنفورد، آزمایشگاهی واقع در شمال کالیفرنیا که پس از معرفی مدل مارک ۱ از ایده‌های روزنبلات استقبال کرده بود، مشغول به کار بود. او در آنجا به‌همراه تیمی بزرگ‌تر از پژوهشگران، در تلاش برای توسعه یک شبکه عصبی با قابلیت خواندن کاراکترهای دست‌نویس، نه‌صرفاً حروف چاپی بود و هدف از ارائه او در کنفرانس، نمایش پیشرفت‌های حاصل در این زمینه تحقیقاتی بود، اما پس از اتمام سخنرانی مونسان و آغاز پاسخگویی به پرسش‌های حاضران، مینسکی با صدای بلند اظهار داشت: «چگونه یک جوان هوشمند همچون شما می‌تواند زمان خود را صرف چنین موضوعی کند؟»

«ران سوگرت» از مهندسان شاغل در آزمایشگاه هوانوردی کرنل (جایی که مارک ۱ متولد شد) که در میان تماشاگران نشسته بود، از این سخنان منجر شد. او از لحن مینسکی به خشم آمد و شک داشت که این حمله ربطی به ارائه جلوی سالن داشته باشد. دغدغه مینسکی تشخیص دست‌نویس نبود، بلکه او به اصل ایده پرسپترون حمله کرد و گفت: «این ایده آینده‌ای ندارد.» در آن‌سوی سالن، «ریچارد دودا» که در تیم سازنده سیستم تشخیص دست‌خط بود، از خنده حضار هنگام تمسخر ادعای تقلید شبکه عصبی مغز توسط

چهارشنبه ۷ خرداد ۱۴۰۴
سال چهارم - شماره ۷۹۷
www.hammihanonline.ir

پرسپترون توسط مینسکی، آزده‌خاطر شد. این رفتار از مینسکی که از بهپا کردن جنجال خوشش می‌آمد، عجیب نبود. او یک‌بار به جمعی از فیزیک‌دانان گفته بود که هوش مصنوعی طی چندسال، بیشتر از فیزیک در قرن‌های پیشرفت کرده است. اما دودا حس می‌کرد که استاد ام‌آی‌تی، دلایل عملی هم برای حمله به کار مراکزی مانند اس‌آرای و کرنل دارد؛ ام‌آی‌تی برای دریافت بودجه تحقیقاتی دولتی با این آزمایشگاه‌ها رقابت می‌کرد. به‌تدر در کنفرانس، وقتی محقق دیگری سیستم جدیدی برای ساخت گرافیک کامپیوتری نشان داد، مینسکی از هوشمندی آن تعریف کرد و یک کنایه دیگر هم به ایده‌های روزنبلات زد: «ایا یک پرسپترون می‌تواند این کار را بکند؟» در پی برگزاری آن کنفرانس، مینسکی و «اسمور پیپرِت»، همکار او در ام‌آی‌تی، کتابی با عنوان «پرسپترون‌ها» در زمینه شبکه‌های عصبی منتشر کردند. بسیاری معتقد بودند که این کتاب، مسیر تحقیقات آینده در مورد ایده‌های روزنبلات را به مدت ۱۵ سال سد کرد. مینسکی و پیپرِت با دقت و ظرافت چشمگیری به‌توصیف پرسپترون پرداختند، به‌طوری‌که در بسیاری از جوانب از شرح خود روزنبلات پیشی گرفت. آنها قابلیت‌های این سیستم، همچنین محدودیت‌های آن را به‌خوبی درک می‌کردند. آنها نشان دادند که پرسپترون قادر به پردازش مفهوم موسوم به «بای انحصاری» (exclusive-or) در اصطلاحات ریاضی نیست؛ مفهومی انتزاعی که دلالت‌های وسیع‌تری را در برداشت.

به‌عنوان مثال، هنگامی که دو نقطه روی یک مربع مقوایی به پرسپترون ارائه می‌شد، این سیستم می‌توانست تشخیص دهد که آیا هر دو نقطه، سیاه یا هر دو، سفید هستند. اما قادر به پاسخگویی به پرسش ساده «آیا این دو نقطه رنگ‌های متفاوتی دارند؟» نبود. این امر نمایانگر آن بود که پرسپترون در برخی موارد حتی از تشخیص الگوهای ساده نیز ناتوان است، چه برسد به الگوهای فوق‌العاده پیچیده‌ای که مشخصه تصاویر هوایی با کلمات گفتاری هستند. برخی محققان از جمله روزنبلات، پیش‌تر در حال بررسی نوع جدیدی از پرسپترون بودند که هدف آن رفع این نقص بود. با این حال، پس از انتشار کتاب مینسکی، تخصیص بودجه‌های دولتی به‌سمت سایر فناوری‌ها معطوف شد و ایده‌های روزنبلات به‌تدریج از عرصه توجه علمی محو شد. در پیروی از رویکرد مینسکی، غالب پژوهشگران به سوی یاداریم موسوم به «هوش مصنوعی نمادین» گرایش یافتند.

هوش مصنوعی نمادین

فرانک روزنبلات در پی طراحی سیستمی بود که قادر به یادگیری رفتار به‌صورت خودکار و مشابه با عملکرد مغز باشد. این رویکرد بعدها توسط دانشمندان تحت‌عنوان «ارتباط‌گرایی» شناخته شد، چراکه همانند مغز، بر مبنای شبکه‌های گسترده از محاسبات متصل به یکدیگر استوار بود. با این وجود، سیستم ابداعی روزنبلات به‌مراتب ساده‌تر از ساختار پیچیده مغز بوده و قابلیت یادگیری آن محدود به موارد نسبتاً ساده بود. مینسکی، همسو با سایر پژوهشگران پیشرو در این حوزه، بر این باور بود که دانشمندان علوم کامپیوتر در بازآفرینی هوش با چالش‌های اساسی روبه‌رو خواهند بود، مگر آن‌که از محدودیت‌های این ایده فراتر رفته و سیستم‌هایی را با رویکردی کاملاً متفاوت و سادتر بنا نهند. در حالی‌که شبکه‌های عصبی از طریق تحلیل داده‌ها به‌صورت خودکار به فراگیری وظایف می‌پردازند، هوش مصنوعی نمادین چنین عملکردی نداشت. این رویکرد براساس دستورالعمل‌های دقیق و مشخصی عمل می‌کرد که توسط مهندسان انسانی تدوین شده بودند؛ مجموعه‌ای از قواعد مجزا که تمام اقدامات موردانتظار از یک ماشین را در هر موقعیت احتمالی تعریف می‌کردند. این رویکرد، هوش مصنوعی نمادین نامیده شد؛ چراکه دستورالعمل‌های مذکور نحوه انجام عملیات خاص روی مجموعه‌های معینی از نمادها، نظیر ارقام و حروف را به ماشین‌ها آموزش می‌دادند.

در طول دهه بعد، این پارادایم بر تحقیقات هوش مصنوعی مسلط شد. این جنبش در اواسط دهه ۱۹۸۰ میلادی با پروژه‌ای به‌نام سایک (Cyc)، تلاشی برای بازآفرینی عقل سلیم از طریق تدوین گام‌به‌گام قواعد منطقی، به اوج آرمان‌های خود رسید. یک تیم کوچک از دانشمندان علوم کامپیوتر، مستقر در آستین، تگزاس، اوقات خود را صرف ثبت حقایق بنیادین نظیر «امکان حضور همزمان در دو مکان وجود ندارد» و «هنگام نوشیدن قهوه از فنجان، باید دهانه باز آن رو به بالا باشد» می‌کردند. آنها آگاه بودند که این فرآیند دهه‌ها و شاید قرن‌ها به‌طول خواهد انجامید، اما همچون بسیاری دیگر، بر این باور بودند که این تنها مسیر ممکن است.

روزنبلات کوشید تا دامنه کاربرد پرسپترون را به فراتر از پردازش تصاویر گسترش دهد. او و همکارانش در دانشگاه کرنل، سیستمی را برای تشخیص کلمات گفتاری طراحی کردند که آن را با الهام از یک داستان کوتاه انگلیسی درباره یک گر به‌شکنگو، «تورموری» نامیدند؛ با این حال، این سیستم هرگز عملکرد مطلوبی نداشت. در اواخر دهه ۱۹۶۰ میلادی، روزنبلات به یک حوزه تحقیقاتی کاملاً متفاوت روی آورد و به انجام آزمایش‌های مغزی روی موش‌های صحرایی پرداخت. پس از آنکه گروهی از موش‌ها توانایی مسیرپایی در یک ماز را آموختند، او بافت مغزی آنها را به گروه دوم تزریق کرد. سپس گروه دوم را در همان ماز قرار داد تا بررسی کند، آیا ذهن آنها اطلاعات آموخته‌نشده توسط گروه اول را جذب کرده یا خیر. نتایج این آزمایش‌ها قطعی و واضح نبود.

در تابستان ۱۹۷۱ میلادی، فرانک روزنبلات در ۴۳سالگی و در سالروز تولدش، در پی یک سانحه قایقرانی در خلیج چسپایک جان باخت. جزئیات دقیق این حادثه در آب توسط روزنامه‌ها منتشر نشد. به گفته یکی از همکارانش، او دو دانشجو را با قایق بادبانی به خلیج برده بود. از آنجا که دانشجویان قایقرانی بلد نبودند، وقتی بوم قایق چرخید و به روزنبلات خورد و او را به آب انداخت، نتوانستند قایق را برگردانند. به این ترتیب، در حالی‌که او در خلیج غرق می‌شد، قایق به راه خود ادامه داد.

مترجم:محسن محمودی

نمایش خانگی

سقوط آرام یک انسان معمولی



در دورانی که رسانه‌های تصویری، در بازتولید روایت‌های مسلط از فرودستان و طبقات پایین جامعه نقش کلیدی دارند، سریال «وحشی» قدیمی متفاوت و جسورانه برمی‌دارد، نه با هدف قهرمان‌سازی از کارگر و نه با نیت سرزنش یا تقبیح او، بلکه با رویکردی فهم‌محور و همدلانه سراغ یکی از آسیب‌پذیرترین لایه‌های جامعه رفته است؛ کارگری که میان اخلاق، قدرت و زخم‌های ساختاری گرفتار شده است. روایت اصلی به‌جای اینکه حول قهرمان یا ضدقهرمانی متعارف بچرخد، روی یک وضعیت انسانی-اجتماعی متمرکز است؛ سقوط تدریجی یک فرد شریف در دل ساختاری ناعادلانه. داوود اشرف، کارگر معنی که در تلاش برای ساختن آینده‌ای بهتر برای خانواده‌اش است، به‌تدریج در مسیر تصمیم‌هایی قرار می‌گیرد که ظاهرشان جنایت است، اما باطن‌شان ترکیبی از فقر، ناتوانی و رنج مزم‌ن است.

سریال با زیرکی، روایت حاشیه‌نشین‌ها را نه از منظر جرم و جنایت، بلکه از زاویه ناامنی اقتصادی، نبود حمایت و انباشت خشم فرودست ارائه می‌دهد. داوود، نه قاتلی بالفطره، که انسانی است آسیب‌پذیر در برابر ساختارهایی که امکان انتخاب‌های اخلاقی را از او سلب کرده‌اند. شخصیت داوود، ترکیبی از تناقض، احساسات و تلاش برای بقاست. او در عین حال که مهربان و خانواده‌دوست است، در بزنگاه‌هایی وارد خشونت می‌شود. این خشونت نه از شرارت، بلکه از بن‌بست و خفگی می‌آید. او کارگری است با آرمان‌هایی کوچک اما شریف، اما ساختارهای اجتماعی-اقتصادی به‌گونه‌ای عمل می‌کنند که حتی همین خواسته‌های حداقلی نیز در دسترس او نیستند. این تضاد، شخصیت داوود را به یک سوژه انسانی و پیچیده تبدیل می‌کند؛ نه معصوم، نه هیولا، بلکه انسانی معمولی در شرایطی غیرمعمول.

سریال موفق می‌شود کارگر را نه‌فقط در جایگاه ابژه‌ای برای ترحم، بلکه به‌مثابه یک سوژه اخلاقی-سیاسی بازنمایی کند که مجبور به انتخاب‌هایی ناانسانی در دل سیستمی ضدانسانی شده است. یکی از نکات کلیدی در تحلیل بازنمایی کارگر در این سریال، درک موقعیت او در ساختار قدرت است. داوود به‌رغم تلاش، هیچ‌گاه از حاشیه به متن قدرت راه نمی‌یابد. او در بهترین حالت، مهره‌ای کوچک در معاملات اقتصادی بزرگ‌تر است. موقعیت داوود به‌طور دائم در لبه‌ی حذف و حذف‌شدگی است. در ساختار رسمی، او نه‌تنها نماینده هیچ قدرتی نیست، بلکه همواره مظنون، تحت‌نظر و فاقد حمایت‌نهادی است. به‌همین دلیل وقتی دست‌به‌خشونت می‌زند، سیستم به‌جای اینکه دلایل اجتماعی را بررسی کند، او را به‌سرعت به‌عنوان مجرم شناسایی می‌کند. سؤال کلیدی‌ای که در دیالوگ معروف «این‌ها نیست اراذل اولیاشه یا کارگره؟!» مطرح می‌شود، به‌خوبی نشان می‌دهد که کارگر، از نگاه قدرت رسمی، فردیتی ندارد، بلکه صرفاً بسته‌ای از برچسب‌هاست؛ مجرم، مظنون، خشن، سریال «وحشی» با ظرافت، داوود را در موقعیت اخلاقی‌ای قرار می‌دهد که در آن نه‌تنها مخاطب، بلکه خود او نیز دچار پرسش و تردید می‌شود. آیا هدف نیک، می‌تواند وسیله نادرست را توجیه کند؟ آیا کسی که تحت فشار دست به خشونت می‌زند، شایسته‌سرزنش است یا ترحم؟ در این موقعیت اخلاقی پیچیده، سریال نه به‌دنبال تطهیر کامل داوود است، نه در پی شیطانی جلوه‌دادن او. او انسانی است که اشتباه می‌کند، اما اشتباهاتش از دل ضعف‌های اخلاقی ذاتی نمی‌آیند، بلکه از محدودیت‌های ساختاری نشأت می‌گیرند. همین امر است که موقعیت اخلاقی داوود را بسیار انسانی، ترازیک و تأمل‌برانگیز می‌سازد. یکی از هوشمندانه‌ترین تصمیم‌های فرمی سریال، انتخاب زاویه دید است. دوربین به‌جای آنکه در جایگاه داور بنشیند، با داوود راه می‌رود، درد می‌کشد، سکوت می‌کند و شکست می‌خورد. این زاویه‌دید، به‌جای آنکه مخاطب را دعوت به‌صدر حکم کند، او را درگیر فهم و درک موقعیت می‌کند. در بسیاری از صحنه‌ها، ما با کلوزآپ‌هایی از صورت گرفته‌ی داوود، دست‌های زبر و نگاه‌های پر از تردید او مواجه‌ایم. هیچ‌گاه او را از فاصله‌ی بالا یا تحقیرآمیز نمی‌بینیم. این هم‌سطحی تصویری، خود یک انتخاب گفتمانی است؛ دیدن کارگر نه به‌مثابه مجرم، بلکه انسانی در حال تلاش برای باقی‌ماندن.

«وحشی» بیش از آنکه روایت فردی یک قتل یا خشونت باشد، بازتابی از وحشی‌گری ساختارهایی است که انسان‌ها را به لبه پرتگاه می‌برند. کارگر در این روایت، نماینده‌ای از میلیون‌ها نفری است که در سکوت، با تورم، فقر، تحقیر و بی‌پناهی دست‌وپنجه‌نرم می‌کنند و او در مواجهه با مرزهای ناپایدی بقا، ناگزیر با انتخاب‌هایی می‌شوند که برای ناظری از بیرون، غیرقابل درک به‌نظر می‌رسد.