

تصحیح روش یادگیری آمیخته تقویت شده با استفاده از آزمون وونگ و کاربرد آن در مدل آمیخته گاما

صدیقه زمانی مهران^۱

تاریخ دریافت: ۱۴۰۱/۰۸/۱۲

تاریخ پذیرش: ۱۴۰۲/۰۱/۱۸

چکیده:

روش یادگیری آمیخته تقویت شده (BML)، روشی فزاینده برای یادگیری مدل‌های آمیخته در مسئله طبقه‌بندی است. در هر مرحله از روش یادگیری آمیخته تقویت شده، مؤلفه جدیدی با توجه به یک تابع هدف در جهت به حداکثر رساندن تابع هدف به مدل آمیخته اضافه می‌شود. از جمله توابع هدف مورد استفاده در این روش، تابع درستنمایی و به‌طور معادل معیارهای اطلاع هستند. در این روش مؤلفه جدیدی به مدل آمیخته اضافه می‌شود که باعث بیشترین افزایش تابع درستنمایی شود. چون تابع درستنمایی و معیارهای اطلاع توانایی تشخیص مدل‌های معادل را ندارد، بنابراین ممکن است مدل آمیخته جدید و مدل آمیخته فعلی معادل باشند و اضافه کردن مؤلفه جدید به مدل آمیخته فعلی باعث بهبود مدل نشود. در این مقاله روش یادگیری آمیخته تقویت شده با استفاده از آزمون انتخاب مدل وونگ که توانایی تشخیص مدل‌های معادل را دارد، تصحیح شده است. همچنین عملکرد دو روش یادگیری با استفاده از داده‌های شبیه‌سازی و مجموعه داده‌های واردات کالای ایالات متحده توسط گمرک ارزیابی شده است.

واژه‌های کلیدی: آزمون وونگ، انتخاب مدل، تابع درستنمایی ماکسیمم، روش یادگیری آمیخته تقویت شده.

۱ مقدمه

مرحله بعد، طبقه‌بندی کننده جدید با طبقه‌بندی کننده فعلی ترکیب می‌شود تا یک طبقه‌بندی پیشرفته بر اساس نظریه تقویت کردن به دست آید که به صورت تحلیلی، کران بالای نرخ خطای آموزشی را به حداقل می‌رساند. این فرآیند برای تولید مجموعه‌ای از طبقه‌بندی کننده‌های پایه تکرار می‌شود. الگوریتم تقویت کردن به‌طور موفقیت آمیزی برای انواع مختلفی از طبقه‌بندی‌ها استفاده می‌شود که می‌تواند دقت را نسبت به طبقه‌بندی کننده‌های پایه بهبود بخشد [۶].

کاربرد الگوریتم تقویت کننده برای برآورد تابع چگالی احتمال به‌طور گسترده مورد مطالعه قرار گرفته است [۱۳]. اخیراً، الگوریتم‌های تقویت کننده در برخی از مسائل یادگیری مدل‌های آمیخته گسترش یافته است [۹] که الگوریتم یادگیری آمیخته تقویت شده (BML) نامیده می‌شود. ایده اصلی الگوریتم یادگیری آمیخته تقویت شده، یادگیری مدل‌های آمیخته به‌طور فزاینده و بازگشتی است. الگوریتم یادگیری آمیخته تقویت شده از یک مدل آمیخته شروع و به تدریج یک مؤلفه آمیخته جدید را به گونه‌ای اضافه می‌کند که همواره تابع هدف از پیش تعریف شده را بهینه کند. روش‌های مشابه دیگری برای یادگیری

مفهوم تقویت کردن به‌طور گسترده در مسائل طبقه‌بندی الگوهای مختلف در یادگیری ماشین استفاده می‌شود. روش تقویت کردن روش کلی برای ترکیب طبقه‌بندی کننده‌های چندگانه است که تقریباً برای هر نوع الگوریتم یادگیری به منظور بهبود دقت طبقه‌بندی کننده‌ها استفاده می‌شود [۱۰]. ایده اصلی روش تقویت کردن این است که به‌طور متوالی مجموعه‌ای از طبقه‌بندی کننده‌های ضعیف را آموزش داده و باهم ترکیب می‌کنیم تا یک طبقه‌بندی قوی و قابل اعتماد ساخته شود [۱۱]. به عنوان مثال، در الگوریتم معروف تقویت کننده تطبیق پذیر [۲]، توزیع احتمالی برای همه نمونه‌های آموزشی در فضای ورودی معرفی می‌شود. در ابتدا، به تمام نمونه‌های آموزشی وزن‌های مساوی اختصاص داده می‌شوند، سپس در تکرارهای آموزشی، وزن آن نمونه‌هایی که طبقه‌بندی آن‌ها سخت‌تر و پیچیده‌تر هستند به تدریج افزایش می‌یابند و به‌طور خودکار این وزن‌ها در تمامی مجموعه آموزشی بر اساس نتایج طبقه‌بندی به‌روزرسانی می‌شوند. با استفاده از وزن‌های آموزشی به‌روزرسانی شده، طبقه‌بندی کننده جدید بر روی این نمونه‌ها آموزش داده می‌شود. در

^۱ گروه آمار، دانشکده علوم پایه، دانشگاه بین‌المللی امام خمینی (ره)، قزوین، ایران (نویسنده مسئول: zamani@sci.ikiu.ac.ir)

را معرفی کرد که در آن K تعداد مدل‌های آمیخته، x بردار ویژگی‌ها، c_k و $f_k(x)$ به ترتیب نشان‌دهنده وزن‌ها و توزیع مؤلفه k -ام مدل آمیخته است.

یادگیری مدل‌های آمیخته به‌طور گسترده در یادگیری ماشین مورد مطالعه قرار گرفته است. پاولویچ [۹] و کیم و پاولویچ [۳] روش یادگیری آمیخته تقویت شده را برای یادگیری مدل‌های آمیخته معرفی کردند. در این روش مؤلفه جدید $f_k(x)$ با وزن c_k به مدل آمیخته $F_{k-1}(x)$ با $k-1$ مؤلفه آمیخته اضافه می‌شود و به مدل آمیخته جدید $F_k(x)$ ،

$$F_k(x) = (1 - c_k)F_{k-1}(x) + c_k f_k(x)$$

با k مؤلفه آمیخته تبدیل می‌شود که در آن c_k نشان‌دهنده وزنی برای ترکیب مؤلفه آمیخته جدید با مدل آمیخته فعلی $F_{k-1}(x)$ است. این روش تا زمانی تکرار می‌شود که برخی از شرایط همگرایی ایجاد شود. هر مؤلفه آمیخته جدید $f_k(x)$ و وزن ترکیب آن c_k را می‌توان بر اساس تابع هدفی مانند $Q(F_k)$ که از پیش تعریف شده، تعیین کرد. الگوریتم یادگیری آمیخته تقویت شده (BML) در جدول ۱ توضیح داده شده است.

جدول ۱ توصیف الگوریتم یادگیری آمیخته تقویت شده

گام اول: برای $k = 1$ مدل اولیه $F_1(x_t)$ را انتخاب کنید.
 گام دوم: برای $k = 2, 3, \dots$ قرار دهید:
 $\{c_k^*, f_k^*\} = \operatorname{argmax}_{c_k, f_k} Q(F_k)$
 اگر مؤلفه‌های بهینه c_k^*, f_k^* یافت شدند مدل آمیخته
 $F_k(x_t) = (1 - c_k^*)F_{k-1}(x_t) + c_k^* f_k^*(x_t)$
 را به‌عنوان مدل آمیخته جدید در نظر گرفته و گام دوم را
 تکرار کنید، در غیر این صورت به گام سوم بروید.
 گام سوم: خروجی مدل آمیخته $F_k(x_t)$ است.

با در نظر گرفتن روش درست‌نمایی ماکسیم (MLE) به‌عنوان روشی برای برآورد پارامترهای مدل آمیخته، تابع هدف را می‌توان تابعی از لگاریتم تابع درست‌نمایی مدل آمیخته $F_k(x)$ تعریف کرد. به عبارت دیگر $Q(F_k) = \sum_{t=1}^T \log F_k(x_t)$ که در آن T تعداد نمونه‌های یادگیری $\{x_1, x_2, \dots, x_T\}$ است.

در گام دوم، مؤلفه جدید $f_k(x)$ با وزن c_k به مدل آمیخته $F_{k-1}(x)$ اضافه می‌شود که ماکسیم‌کننده تابع هدف $Q(F_k)$ باشد. به عبارت دیگر

$$Q((1 - \epsilon)F_{k-1} + \epsilon f_k) > Q(F_{k-1}) \quad (1)$$

مدل‌های آمیخته وجود دارد [۵، ۱۴]، اما نکته اساسی روش یادگیری آمیخته تقویت شده این است که مؤلفه آمیخته جدید در هر مرحله با توجه به تابع هدف برآورد می‌شود، به‌طوری‌که هر مؤلفه جدید همیشه وقتی اضافه می‌شود که تابع هدف را بیش از همه افزایش می‌دهد. نیا و همکاران [۸] مدل تقویت گرادیان با مدل آمیخته گاوسی را معرفی کردند و به پیش‌بینی جریان ماهانه با استفاده از این مدل پرداختند. دوو و همکاران [۱] تابع لگاریتم درست‌نمایی مدل آمیخته را به‌عنوان تابع هدف در نظر گرفتند و در هر تکرار مؤلفه جدیدی را به مدل آمیخته اضافه می‌کردند به‌طوری‌که تابع لگاریتم درست‌نمایی مدل آمیخته جدید نسبت به تابع لگاریتم درست‌نمایی مدل آمیخته فعلی افزایش یابد. در بعضی از موارد ممکن است این افزایش ناچیز باشد و مدل آمیخته جدید و مدل آمیخته فعلی معادل باشند. به عبارت دیگر روش یادگیری آمیخته تقویت شده توانایی تشخیص مدل‌های معادل را ندارد. در این مقاله، روش یادگیری آمیخته تقویت شده بر اساس آزمون انتخاب مدل وونگ تصحیح شده است. اشنایدر و همکاران [۱۲] برای مقایسه مدل‌های تک متغیره و چندمتغیره آشیانی از آزمون وونگ استفاده کردند و با استفاده از مطالعات شبیه‌سازی و داده‌های واقعی نشان دادند وقتی که مدل‌ها آشیانی هستند این آزمون معمولاً به‌خوبی یا گاهی بهتر از آزمون نسبت درست‌نمایی عمل می‌کند.

مقاله به شرح زیر سازمان‌دهی شده است: در بخش ۲، به بررسی الگوریتم یادگیری آمیخته تقویت شده برای مدل‌های آمیخته پرداخته شده است. در بخش ۳، الگوریتم یادگیری آمیخته تقویت شده برای مدل‌های آمیخته بر اساس آزمون انتخاب مدل تصحیح شده است و مثالی از مدل آمیخته گاما ارائه شده است. در بخش ۴، عملکرد دو روش یادگیری آمیخته تقویتی تصحیح شده و یادگیری آمیخته تقویت شده با استفاده از مطالعات شبیه‌سازی مقایسه شده است. آنالیز داده‌های واقعی واردات کالاهای ایالات متحده توسط گمرک در بخش ۵ ارائه شده است.

۲ مروری بر الگوریتم یادگیری آمیخته تقویت شده

در این بخش، به‌مرور الگوریتم یادگیری آمیخته تقویت شده (BML) پرداخته که کاربرد فراوانی در بسیاری از مسائل طبقه‌بندی الگوها دارد. مک لاکلان [۶]، مدل آمیخته $F_K(x)$

$$F_K(x) = \sum_{k=1}^K c_k f_k(x), c_k \geq 0, \sum_{k=1}^K c_k = 1$$

هدف، تابع لگاریتم مدل آمیخته است. به عبارت دیگر تابع $f_k^*(x)$ با وزن c_k^* به مدل $F_{k-1}(x)$ اضافه می‌شود اگر $Q(F_k) > Q(F_{k-1})$ باشد به طوری که $F_k(x_t) = (1 - c_k^*)F_{k-1}(x_t) + c_k^*f_k^*(x_t)$ است. توجه کنید که افزایش اندک مقدار $Q(F_k)$ نسبت به $Q(F_{k-1})$ باعث انتخاب مدل آمیخته $F_k(x_t)$ به عنوان مدل بهینه می‌شود، اما ممکن است که دو مدل $F_k(x_t)$ و $F_{k-1}(x_t)$ معادل باشند. به عبارت دیگر نامساوی معرفی شده توسط دوو و همکاران [۱] توانایی تشخیص مدل‌های معادل را ندارد. اگر دو مدل $F_k(x_t)$ و $F_{k-1}(x_t)$ معادل باشند مدل $F_{k-1}(x_t)$ به عنوان مدل بهینه انتخاب می‌شود، در صورتی که بر اساس نامساوی معرفی شده توسط دوو و همکاران [۱] $Q((1 - \epsilon)F_{k-1} + \epsilon f_k) > Q(F_{k-1})$ با وزن c_k^* به مدل آمیخته $F_{k-1}(x_t)$ اضافه می‌شود. در این مقاله به تصحیح الگوریتم معرفی شده توسط دوو و همکاران [۱] بر اساس آزمون انتخاب مدل که توانایی تشخیص مدل‌های معادل را دارند، پرداخته شده است.

۳ تصحیح الگوریتم یادگیری آمیخته تقویت شده

معمولاً توزیع درست پدیده‌های تحت مطالعه معلوم نیست. به عبارت دیگر نمونه تصادفی $X_1, \dots, X_T$ از توزیع مجهول $h(x_t)$ پیروی می‌کند. برای انجام هر استنباط آماری نیاز به توزیع درست جامعه داریم. از دیدگاه ریاضی $h(x_t)$ ترکیبی از منحنی‌ها است که نقاط مشاهده شده را به هم وصل می‌کند. به دلیل پیچیدگی این نوع مدل‌ها، فرض می‌شود که نمونه تصادفی از یک توزیع (مدل) پارامتری پیروی می‌کند. مطالعه و بررسی مدل یا ساختار داده‌های مشاهده شده، انتخاب مدل نامیده می‌شود. برای این منظور خانواده‌ای از مدل‌ها مانند $F = \{f^{\theta_i}(x_t), i = 1, 2, \dots, k\}$ در نظر گرفته می‌شود و از بین مدل‌های پیشنهادی مدل بهینه تحت تابع هدف خاصی انتخاب می‌شود.

بردار تصادفی $X(T \times 1)$ و فرضیه‌های

$$H_1: X \sim f^{\theta_k}(x_t)$$

و

$$H_0: X \sim f^{\theta_{k-1}}(x_t)$$

را در نظر بگیرید که در آن $f^{\theta_k}(x_t)$ مدل آمیخته $F_k(x_t)$ و $f^{\theta_{k-1}}(x_t)$ مدل آمیخته $F_{k-1}(x_t)$ هستند. توجه کنید که دو مدل پیشنهادی (رقیب) $F_k(x_t)$ و $F_{k-1}(x_t)$ مدل‌های آشیانی هستند. به این مفهوم که

که در آن ϵ نشان دهنده انحراف کوچک ثابت است. با استفاده از بسط تیلور مرتبه اول، می‌توان بسط جمله سمت چپ معادله ۱ را به صورت

$$\begin{aligned} Q((1 - \epsilon)F_{k-1} + \epsilon f_k) &= Q(F_{k-1} + \epsilon(f_k - F_{k-1})) \\ &= Q(F_{k-1}) + \epsilon \langle \nabla Q(F_{k-1}), (f_k - F_{k-1}) \rangle \\ &\quad + O(\epsilon(f_k - F_{k-1})) \\ &\approx Q(F_{k-1}) + \epsilon \langle \nabla Q(F_{k-1}), (f_k - F_{k-1}) \rangle \end{aligned} \quad (2)$$

نوشت که در آن

$$\nabla Q(F_{k-1}) = \nabla Q(F)|_{F=F_{k-1}} = \frac{1}{F_{k-1}(x_t)}$$

و $\langle P_1, P_2 \rangle$ نشان دهنده ضرب داخلی تابعی برای دو مدل آمیخته $P_1(x_t)$ و $P_2(x_t)$ است. به عبارت دیگر

$$\langle P_1, P_2 \rangle = \frac{1}{T} \sum_{t=1}^T P_1(x_t) P_2(x_t).$$

اگر ϵ به اندازه کافی کوچک باشد، می‌توان همه جملات با مرتبه بالا را نادیده گرفت؛ بنابراین، تابع هدف بر اساس مدل آمیخته جدید را می‌توان با عبارات خطی ۲ تقریب زد، به دوو و همکاران [۱] مراجعه کنید.

مؤلفه جدید f_k باید به گونه‌ای برآورد شود که تابع هدف

$$Q((1 - \epsilon)F_{k-1} + \epsilon f_k)$$

را به سرعت افزایش دهد. با اعمال تقریب خطی ۲ در معادله ۱، انتخاب مؤلفه جدید به مفهوم انتخاب f_k ای است که ماکسیمم کننده $\langle \nabla Q(F_{k-1}), (f_k - F_{k-1}) \rangle$ باشد. به عبارت دیگر

$$\begin{aligned} f_k^* &= \operatorname{argmax}_{f_k} \langle \nabla Q(F_{k-1}), (f_k - F_{k-1}) \rangle \\ &= \operatorname{argmax}_{f_k} \frac{1}{T} \sum_{t=1}^T \frac{(f_k(x_t) - F_{k-1}(x_t))}{F_{k-1}(x_t)} \\ &= \operatorname{argmax}_{f_k} \frac{1}{T} \sum_{t=1}^T \frac{f_k(x_t)}{F_{k-1}(x_t)}. \end{aligned} \quad (3)$$

در معادله ۳ انتخاب مؤلفه جدید $f_k(x_t)$ بر اساس روش درست‌نمایی ماکسیمم (ML) است. به شرط برآورد $f_k(x)$ ، وزن c_k^* به صورت

$$c_k^* = \operatorname{argmax}_{c_k \in [0, 1]} Q((1 - c_k)F_{k-1} + c_k f_k^*) \quad (4)$$

قابل محاسبه است. در عمل، وزن بهینه c_k^* را می‌توان با استفاده از جستجوی شبکه‌ای در بازه‌ی $[0, 1]$ پیدا کرد.

دوو و همکاران [۱] برای تعیین k و $f_k^*(x)$ و وزن بهینه c_k^* ، نامساوی $Q((1 - \epsilon)F_{k-1} + \epsilon f_k) > Q(F_{k-1})$ را معرفی کردند که در آن تابع

باشد فرضیه معادل بودن دو مدل را می‌پذیریم. با توجه به این‌که مدل‌های پیشنهادی آشیانی هستند، بنابراین پذیرش فرضیه معادل بودن دو مدل منتج به پذیرش مدل $f^{\theta_{k-1}}(x_t)$ می‌شود. اگر مقدار آماره $2LR(\hat{\theta}_k, \hat{\theta}_{k-1})$ از مقدار بحرانی $\chi^2_{1-\frac{\alpha}{2}}(1)$ بیشتر باشد، فرضیه صفر را رد می‌کنیم به طوری‌که مدل $f^{\theta_k}(x_t)$ برای برازش بر روی داده‌ها مناسب‌تر از $f^{\theta_{k-1}}(x_t)$ است. اگر مقدار آماره $2LR(\hat{\theta}_k, \hat{\theta}_{k-1})$ از مقدار بحرانی $\chi^2_{1-\frac{\alpha}{2}}(1)$ کمتر باشد، مدل $f^{\theta_{k-1}}(x_t)$ برای برازش بر روی داده‌ها مناسب‌تر از $f^{\theta_k}(x_t)$ است. با توجه به نواحی رد و پذیرش می‌توان نتیجه گرفت وقتی که $\chi^2_{1-\frac{\alpha}{2}}(1) > 2LR(\hat{\theta}_k, \hat{\theta}_{k-1})$ باشد می‌توان مدل $f^{\theta_k}(x_t)$ را برای برازش بر روی داده‌ها انتخاب کرد. به عبارت دیگر می‌توان مؤلفه جدید $f_k(x)$ با وزن c_k به مدل آمیخته $F_{k-1}(x)$ اضافه کرد؛ بنابراین اگر تابع لگاریتم درستنمایی مدل آمیخته به عنوان تابع هدف در نظر گرفته شود، هنگامی مؤلفه جدید $f_k(x)$ به مدل آمیخته فعلی اضافه می‌شود که

$$Q((1-\epsilon)F_{k-1} + \epsilon f_k) - Q(F_{k-1}) > \frac{1}{2} \chi^2_{1-\frac{\alpha}{2}}(1) \quad (5)$$

الگوریتم یادگیری آمیخته تقویتی تصحیح شده (CBML) در جدول ۲ توضیح داده شده است.

جدول ۲. توصیف الگوریتم یادگیری آمیخته تقویتی تصحیح شده

گام اول: برای $k=1$ مدل اولیه $F_1(x_t)$ را بر اساس آزمون وونگ انتخاب کنید.

گام دوم: برای $k=2, 3, \dots$ قرار دهید:

$$\{c_k^*, f_k^*\} = \operatorname{argmax}_{c_k, f_k} Q(F_k)$$

اگر مؤلفه‌های بهینه c_k^*, f_k^* یافت شدند مدل آمیخته

$$F_k(x_t) = (1 - c_k^*)F_{k-1}(x_t) + c_k^* f_k^*(x_t)$$

را به عنوان مدل آمیخته جدید در نظر گرفته و اگر

$$Q((1-\epsilon)F_{k-1} + \epsilon f_k) - Q(F_{k-1}) > \frac{1}{2} \chi^2_{1-\frac{\alpha}{2}}(1)$$

باشد گام دوم را تکرار کنید، در غیر این صورت به گام سوم بروید.

گام سوم: خروجی مدل آمیخته $F_k(x_t)$ است.

در گام دوم روش یادگیری آمیخته تقویتی تصحیح شده، برای محاسبه مؤلفه‌های بهینه c_k^* و f_k^* باید تابع لگاریتم درستنمایی مدل آمیخته را ماکسیم کرد. چون برآوردگر ماکسیم درستنمایی پارامترهای مدل آمیخته جواب صریحی ندارد، بنابراین روش‌های دیگر از جمله الگوریتم

مدل $F_{k-1}(x_t)$ با ایجاد محدودیت‌های پارامتری بر روی مدل $F_k(x_t)$ ایجاد می‌شود. همچنین فرض کنید $l(\hat{\theta}_i)$ تابع لگاریتم درستنمایی مدل ارائه شده تحت فرضیه‌های $H_i, i=1, 2$ باشد. هدف از انتخاب مدل، انتخاب مدل آمیخته $F_k(x_t)$ یا مدل آمیخته $F_{k-1}(x_t)$ یا هر دو مدل به عنوان مدل‌های معادل است. با در نظر گرفتن دو مدل آمیخته $F_k(x_t)$ و $F_{k-1}(x_t)$ به عنوان دو مدل پیشنهادی (رقیب) آشیانی، مدلی به عنوان مدل بهینه انتخاب می‌شود که به مدل درست نزدیک‌تر باشد. به عبارت دیگر مدلی را انتخاب می‌کنیم که حداقل ریسک کولبک - لیبلر [۴] را داشته باشد. وونگ [۱۵] آزمون فرضیه

$$H_0 : E_h \left(\log \frac{f^{\theta_k}(x_t)}{f^{\theta_{k-1}}(x_t)} \right) = 0$$

در مقابل

$$H_f : E_h \left(\log \frac{f^{\theta_k}(x_t)}{f^{\theta_{k-1}}(x_t)} \right) > 0$$

یا

$$H_g : E_h \left(\log \frac{f^{\theta_k}(x_t)}{f^{\theta_{k-1}}(x_t)} \right) < 0$$

را برای انتخاب مدل در بین مدل‌های آشیانی پیشنهادی بد-توصیف شده، معرفی کرد که مبتنی بر آماره نسبت درستنمایی است. پذیرش H_0 به مفهوم معادل بودن دو مدل $f^{\theta_k}(x_t)$ و $f^{\theta_{k-1}}(x_t)$ است. پذیرش H_f به این مفهوم است که مدل $f^{\theta_k}(x_t)$ برای برازش بر روی داده‌ها مناسب‌تر از $f^{\theta_{k-1}}(x_t)$ است، همچنین پذیرش H_g به این مفهوم است که مدل $f^{\theta_{k-1}}(x_t)$ برای برازش بر روی داده‌ها مناسب‌تر از $f^{\theta_k}(x_t)$ است. این آزمون، زمانی که مدل خوب-توصیف شده باشد تبدیل به آزمون نسبت درستنمایی نیم - پیروسون می‌شود. تعریف می‌کنیم

$$LR(\hat{\theta}_k, \hat{\theta}_{k-1}) = l(\hat{\theta}_k) - l(\hat{\theta}_{k-1})$$

$$= \sum_{t=1}^T \log f^{\hat{\theta}_k}(x_t) - \sum_{t=1}^T \log f^{\hat{\theta}_{k-1}}(x_t)$$

$$= \sum_{t=1}^T \log \left(\frac{f^{\hat{\theta}_k}(x_t)}{f^{\hat{\theta}_{k-1}}(x_t)} \right).$$

وونگ [۱۵] نشان داد که تحت فرضیات $A_1 - A_6$ (ارائه شده در وونگ [۱۵]) و تحت فرضیه H_0 آماره $2LR(\hat{\theta}_k, \hat{\theta}_{k-1})$ دارای توزیع حدی مجموع وزنی توزیع کای اسکور است. اگر مدل‌های رقیب خوب-توصیف شده باشند این توزیع حدی تبدیل به توزیع کای اسکور می‌شود. به عبارت دیگر

$$2LR(\hat{\theta}_k, \hat{\theta}_{k-1}) \xrightarrow{d} \chi^2(1).$$

با انتخاب سطح معنی‌داری α ، اگر

$$\chi^2_{1-\frac{\alpha}{2}}(1) < 2LR(\hat{\theta}_k, \hat{\theta}_{k-1}) < \chi^2_{1-\frac{\alpha}{2}}(1)$$

برآورد می‌شود.

در گام- M ، ابتدا $\lambda_j^{(i+1)}$ و $a_j^{(i+1)}$ با ماکسیم کردن $E(Q_c(F_k)|X, \theta^{(i)})$ به ترتیب نسبت به λ_j و a_j محاسبه می‌شود. برای این منظور از $E(Q_c(F_k)|X, \theta^{(i)})$ نسبت به λ_j و a_j مشتق گرفته و معادل با صفر قرار می‌دهیم. با حل این معادله داریم

$$\lambda_j^{(i+1)} = \frac{1}{T} \sum_{t=1}^T z_{tj}^{(i)} \quad (7)$$

و

$$\frac{E(Q_c(F_k)|X, \theta^{(i)})}{\partial a_j} = \sum_{t=1}^T z_{tj}^{(i)} [\log(x_t) + \log(B_j^{(i)}) - D(a_j)] = 0 \quad (8)$$

که در آن $D(a_j)$ نشان‌دهنده تابع دی-گاما است. سپس $B_j^{(i+1)}$ را با ماکسیم کردن $E(Q_c(F_k)|X, \theta^{(i)})$ نسبت به B_j اما در مقدار $a_j^{(i+1)}$ محاسبه می‌کنیم. به عبارت دیگر

$$B_j^{(i+1)} = \frac{a_j^{(i+1)} \sum_{t=1}^T z_{tj}^{(i)}}{\sum_{t=1}^T x_t z_{tj}^{(i)}}.$$

۴ مطالعه شبیه‌سازی

در این بخش عملکرد دو الگوریتم یادگیری ذکر شده با استفاده از مطالعات شبیه‌سازی مقایسه می‌شود. مدل آمیخته گاما به عنوان مدل درست در نظر گرفته شده و نمونه‌ای به اندازه $T = 400$ از مدل درست تولید شده که در آن

$$k = 3, a_1 = 20, a_2 = 32, a_3 = 40, B_1 = 14, B_2 = 10, B_3 = 6,$$

$$c_1 = 0.5, c_2 = 0.3, c_3 = 0.2$$

است. هدف برازش مدل آمیخته گاما با استفاده از دو روش یادگیری آمیخته تقویت شده و یادگیری آمیخته تقویتی تصحیح شده بر اساس تابع درستنمایی ماکسیم است. مقادیر برآورد پارامترهای مدل و تابع هدف تحت دو روش یادگیری آمیخته به ترتیب در جدول‌های ۳ و ۴ ارائه شده است. در جدول ۳، مقادیر برآورد پارامترهای مدل و تابع هدف تحت روش یادگیری آمیخته تقویتی تصحیح شده گزارش شده است. با توجه به جدول ۳ مشاهده می‌شود که مدل برآورد شده آمیخته با ۳ مؤلفه آمیخته

$$F_3(x_t) = 0.498 \cdot G(237.1348, 15.9789) + 0.2493 \cdot G(32.0999, 7.1645) + 0.2527 \cdot G(35.0403, 4.7510)$$

مدلی مناسب برای برآورد بر روی داده‌های تولید شده است. با در نظر گرفتن $\alpha = 0.05$ ، از جدول ۳ مشاهده می‌شود چون مقدار آماره $2LR(\hat{\theta}_k, \hat{\theta}_{k-1})$ برای $k = 2, 3$ بزرگ‌تر از $\chi^2_{0.95}(1) = 5.0238$ است، بنابراین مؤلفه‌های $f_3^{\Phi_2}(x_t)$ و $f_3^{\Phi_3}(x_t)$ به مدل اضافه می‌شود اما

EM را می‌توان برای برآورد پارامترهای مدل استفاده کرد. برای این منظور در ادامه مدل آمیخته گاما به عنوان مثالی از مدل‌های آمیخته بیان می‌شود و پارامترهای مجهول مدل با استفاده از الگوریتم EM برآورد می‌شوند.

مدل آمیخته گاما [۱۶] را در نظر بگیرید که هر مؤلفه آن از توزیع گاما با تابع چگالی

$$f_k^{\Phi_k}(x_t) = \frac{1}{B^{a_k} \Gamma(a_k)} x_t^{a_k} e^{-\frac{x_t}{B_k}}$$

پیروی می‌کند که در آن $\Gamma(\cdot)$ تابع گاما است و $\Phi_k = (a_k, B_k)$ نشان‌دهنده پارامترهای k -امین توزیع گاما است. همچنین پارامترهای مدل آمیخته $F_k(x_t)$ را با نماد $\theta_k = (\Phi_k, \theta_{k-1})$ نمایش می‌دهیم. با توجه به معادلات ۳ و ۴ داریم

$$\{f_k^*, c_k^*\} = \operatorname{argmax}_{c_k, f_k} Q((1 - c_k)F_{k-1} + c_k f_k).$$

چون برآوردگر ماکسیم درستنمایی پارامترهای مدل آمیخته جواب صریحی ندارد، بنابراین پارامترهای مدل را با استفاده از الگوریتم EM برآورد می‌کنیم. فرض کنید $X = (x_1, \dots, x_T)'$ نمونه‌ای تصادفی به اندازه T از توزیع آمیخته گاما

$$F_k(x_t) = \sum_{j=1}^k \lambda_j f_j^{\Phi_j}(x_t)$$

با تابع لگاریتم درستنمایی

$$Q(F_k) = \sum_{t=1}^T \log \sum_{j=1}^k \lambda_j f_j^{\Phi_j}(x_t)$$

باشد. چون ماکسیم کردن تابع $Q(F_k)$ سخت و پیچیده هست، بنابراین متغیر نشان‌گر Z_{tj} را به منظور استفاده از الگوریتم EM تعریف می‌کنیم. اگر مشاهده t -ام متعلق به مؤلفه j -ام باشد آنگاه $Z_{tj} = 1$ ، در غیر این صورت $Z_{tj} = 0$ است. تابع لگاریتم درستنمایی کامل بر اساس داده‌های کامل $X_c = (X, Z)$ به صورت

$$Q_c(F_k) = \sum_{t=1}^T \sum_{j=1}^k Z_{tj} \log \{ \lambda_j f_j^{\Phi_j}(x_t) \}$$

قابل محاسبه هست. برای الگوریتم EM ، به طور مکرر تابع امید ریاضی شرطی $E(Q_c(F_k)|X, \theta^{(i)})$ را ماکسیم می‌کنیم که در آن $\theta^{(i)}$ مقدار فعلی در i -امین تکرار و $\theta^{(e)}$ نقطه شروع مشخص شده توسط کاربر است.

در گام- E ، برای i -امین تکرار $Z_{tj}^{(i)}, i = 1, 2, \dots$ به صورت

$$Z_{tj}^{(i)} = P_{\theta^{(i)}}(Z_{tj} = 1 | x_t) = \frac{\lambda_j f_j^{\Phi_j}(x_t)}{F_j(x_t)} \quad (6)$$

و این اختلاف باعث می شود که مدل $f_{\Psi}^{\Phi}(x_t)$ بر اساس روش آمیخته تقویت شده به مدل آمیخته فعلی اضافه شود، اما این اختلاف بسیار ناچیز است؛ بنابراین مدل درست تحت روش یادگیری آمیخته تقویت شده به عنوان مدل بهینه انتخاب نشده است، به عبارت دیگر مدل آمیخته ای با تعداد مؤلفه های آمیخته بیشتر به عنوان مدل بهینه انتخاب شده است که برای برازش این مدل نیاز به برآورد تعداد بیشتری از مؤلفه های آمیخته (پارامتر) است.

از آنجایی که مقدار آماره $2LR(\hat{\theta}_k, \hat{\theta}_{k-1})$ برای $k = 4$ کوچک تر از $\chi_{0.975}^2(1) = 5.0238$ است بنابراین دو مدل آمیخته $F_2(x_t)$ و $F_4(x_t)$ مدل های معادل هستند، پس مؤلفه $f_{\Psi}^{\Phi}(x_t)$ به مدل اضافه نمی شود. مقادیر برآورد پارامترهای مدل و تابع هدف تحت روش یادگیری آمیخته تقویت شده در جدول ۴ ارائه شده است. با توجه به جدول ۴ مشاهده می شود که تحت روش یادگیری آمیخته تقویت شده، مدل آمیخته با ۵ مؤلفه آمیخته به عنوان مدل مناسب برای برازش بر روی داده ها انتخاب شده است. وقتی که $K = 4$ است اگرچه اختلاف تابع لگاریتم درستنمایی دو مدل آمیخته، 0.0073 ، بزرگ تر از 0.0000 است

جدول ۳. مقادیر برآورد پارامترهای مدل و تابع هدف تحت روش یادگیری آمیخته تقویتی تصحیح شده

	$K = 1$	$K = 2$	$K = 3$	$K = 4$
$\hat{c}_1$	۱/۰۰۰۰	۰/۶۶۶۴	۰/۴۹۸۰	۰/۴۹۷۹
$\hat{a}_1$	۲۳/۱۳۴۸	۲۳/۱۳۴۸	۲۳/۱۳۴۸	۲۳/۱۳۴۸
$\hat{B}_1$	۱۵۹۷۸۹	۱۵۹۷۸۹	۱۵۹۷۸۹	۱۵۹۷۸۹
$\hat{c}_2$		۰/۳۳۳۶	۰/۲۴۹۳	۰/۲۴۹۳
$\hat{a}_2$		۳۲/۲۰۹۹	۳۲/۲۰۹۹	۳۲/۲۰۹۹
$\hat{B}_2$		۷/۱۶۴۵	۷/۱۶۴۵	۷/۱۶۴۵
$\hat{c}_3$			۰/۲۵۲۷	۰/۲۵۲۶
$\hat{a}_3$			۳۵/۰۴۰۳	۳۵/۰۴۰۳
$\hat{B}_3$			۴/۷۵۱۰	۴/۷۵۱۰
$\hat{c}_4$				۰/۰۰۰۲
$\hat{a}_4$				۲۲/۶۱۶۷
$\hat{B}_4$				۷/۶۰۹۸
$2LR(\hat{\theta}_k, \hat{\theta}_{k-1})$		۲۰۰/۶۹۳۹	۴۷/۸۹۷۹	۰/۰۱۴۶

موجود است. نمودار تابع چگالی و نمودار چندک - چندک این داده ها در شکل ۲ نمایش داده شده است. شکل ۲ نشان می دهد که در نمودار چندک-چندک، تعداد زیادی از مشاهدات نزدیک به خط ۴۵ درجه (نیمساز اول و سوم) نیستند، به عبارت دیگر مشاهدات از توزیع نرمال پیروی نمی کنند. همچنین با توجه به مقادیر ضریب چولگی و ضریب کشیدگی این مشاهدات (که به ترتیب برابر 0.7360 و -1841.01 هستند) و p -مقدار آزمون شاپیرو که کوچک تر از 0.05 است این نتیجه نیز تأیید می شود، پس فرض نرمال بودن مشاهدات رد می شود. همچنین نمودار تابع چگالی، شکل ۲، دو مدی بودن داده ها را نشان می دهد، بنابراین مدل آمیخته گاما برای برازش بر روی این مجموعه از داده ها پیشنهاد و پارامترهای مدل با استفاده از روش یادگیری آمیخته تقویتی تصحیح شده برآورد شده است.

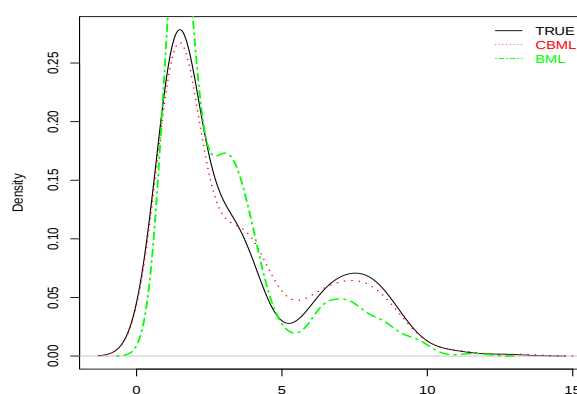
نمودار تابع چگالی داده های تولید شده و مدل های برازش داده شده بر روی این داده ها با استفاده از دو روش یادگیری ذکر شده در شکل ۱ نمایش داده شده است. این نمودار نشان می دهد که مدل برآورد شده با استفاده از روش یادگیری آمیخته تقویتی تصحیح شده برازش مناسب تری را ایجاد می کند.

۵ آنالیز داده های حقیقی

مجموعه داده های مربوط به واردات کالاهای ایالات متحده توسط گمرک (برحسب میلیون دلار) طی سال ها ۱۹۸۹ تا ۲۰۲۲ را در نظر بگیرید. این داده ها در

جدول ۴. مقادیر برآورد پارامترهای مدل و تابع هدف تحت روش یادگیری آمیخته تقویت شده

	$K=1$	$K=2$	$K=3$	$K=4$	$K=5$	$K=6$
$\hat{c}_1$	۱/۰۰۰۰	۰/۶۶۶۴	۰/۴۹۸۰	۰/۴۹۷۹	۰/۴۸۷۶	۰/۴۸۶۹
$\hat{a}_1$	۲۳/۱۳۴۸	۲۳/۱۳۴۸	۲۳/۱۳۴۸	۲۳/۱۳۴۸	۲۳/۱۳۴۸	۲۳/۱۳۴۸
$\hat{B}_1$	۱۵/۹۷۸۹	۱۵/۹۷۸۹	۱۵/۹۷۸۹	۱۵/۹۷۸۹	۱۵/۹۷۸۹	۱۵/۹۷۸۹
$\hat{c}_2$		۰/۳۳۳۶	۰/۲۴۹۳	۰/۲۴۹۳	۰/۲۴۴۲	۰/۲۴۳۷
$\hat{a}_2$		۳۲/۲۰۹۹	۳۲/۲۰۹۹	۳۲/۲۰۹۹	۳۲/۲۰۹۹	۳۲/۲۰۹۹
$\hat{B}_2$		۷/۱۶۴۵	۷/۱۶۴۵	۷/۱۶۴۵	۷/۱۶۴۵	۷/۱۶۴۵
$\hat{c}_3$			۰/۲۵۲۷	۰/۲۵۲۶	۰/۲۴۷۴	۰/۲۴۷۰
$\hat{a}_3$			۳۵/۰۴۰۳	۳۵/۰۴۰۳	۳۵/۰۴۰۳	۳۵/۰۴۰۳
$\hat{B}_3$			۴/۷۵۱۰	۴/۷۵۱۰	۴/۷۵۱۰	۴/۷۵۱۰
$\hat{c}_4$				۰/۰۰۰۲	۰/۰۰۰۲	۰/۰۰۰۱
$\hat{a}_4$				۲۲/۶۱۶۷	۲۲/۶۱۶۷	۲۲/۶۱۶۷
$\hat{B}_4$				۷/۶۰۹۸	۷/۶۰۹۸	۷/۶۰۹۸
$\hat{c}_5$					۰/۰۲۰۶	۰/۰۲۰۶
$\hat{a}_5$					۲۴/۷۰۹۸	۲۴/۷۰۹۸
$\hat{B}_5$					۷/۵۱۱۸	۷/۵۱۱۸
$\hat{c}_6$						۰/۰۰۱۷
$\hat{a}_6$						۲۲/۶۶۳۸
$\hat{B}_6$						۱۶/۰۱۶۶
$LR(\hat{\theta}_k, \hat{\theta}_{k-1})$		۱۰۰/۳۴۶۹	۲۳/۹۴۸۶	۰/۰۰۷۳	۰/۰۲۲۷	-۱/۶۲۴۴



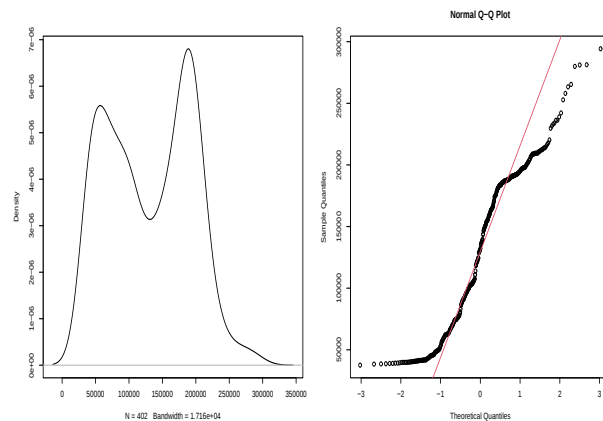
شکل ۱: نمودار تابع چگالی داده‌های تولید شده و مدل‌های برازش داده شده.

مقدار آماره $2LR(\hat{\theta}_k, \hat{\theta}_{k-1})$ برای $k=2, 3$ با $\chi^2_{0.975}(1)$ از جدول ۵ مشاهده می‌شود که مدل آمیخته برآورد شده

$$F_7(x_t) = 0.4800G(6/1044, 1276284323) + 0.5200G(433331, 43457589)$$

پارامترهای مدل پیشنهادی با استفاده از روش یادگیری آمیخته تقویتی تصحیح شده محاسبه و مقادیر پارامترهای برآورد شده و تابع هدف در جدول ۵ ارائه شده است. با در نظر گرفتن $\alpha = 0.05$ و مقایسه

مدلی مناسب برای برآورد بر روی این مجموعه از داده‌ها است. برای $k = 2$ مقدار آماره $2LR(\hat{\theta}_k, \hat{\theta}_{k-1}) = 1409742$ بزرگ‌تر از $\chi^2_{0.975}(1) = 50238$ است بنابراین مؤلفه‌های $f_2^{\Phi_2}(x_t)$ به مدل اضافه می‌شود، اما از آنجایی که مقدار آماره $2LR(\hat{\theta}_k, \hat{\theta}_{k-1}) = 16849$ بزرگ‌تر از $\chi^2_{0.975}(1) = 50238$ است بنابراین دو مدل آمیخته



شکل ۲: نمودار تابع چگالی و نمودار چندک-چندک مشاهدات.

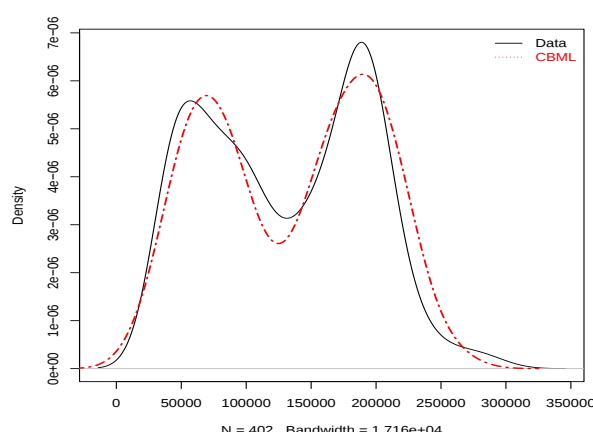
جدول ۵. مقادیر برآورد پارامترهای مدل و تابع هدف تحت روش یادگیری آمیخته تقویتی تصحیح شده

	$K = 1$	$K = 2$	$K = 3$
$\hat{c}_1$	1,0000	0,4800	0,4793
$\hat{a}_1$	61044	61044	61044
$\hat{B}_1$	12762,84323	12762,84323	12762,84323
$\hat{c}_2$		0,5200	0,5193
$\hat{a}_2$		43,3331	43,3331
$\hat{B}_2$		4345,7589	4345,7589
$\hat{c}_3$			0,0013
$\hat{a}_3$			503,827
$\hat{B}_3$			381,7906
$2LR(\hat{\theta}_k, \hat{\theta}_{k-1})$		9742,140	6849,1

لگاریتم درستنمایی (معیار اطلاع) مدل آمیخته فعلی افزایش (کاهش) می‌یابد. در بعضی از موارد ممکن است این افزایش ناچیز باشد و مدل آمیخته جدید و مدل آمیخته فعلی معادل باشند، به این مفهوم که اضافه کردن مؤلفه جدید باعث بهبود مدل فعلی نمی‌شود. به عبارت دیگر در این روش، اختلاف تابع درستنمایی (معیار اطلاع) دو مدل آمیخته فعلی و آمیخته جدید با صفر مقایسه می‌شود. مؤلفه جدید در صورتی به مدل

۶ بحث و نتیجه‌گیری

تابع لگاریتم درستنمایی و معیارهای اطلاع از جمله توابع هدفی هستند که در روش یادگیری آمیخته تقویت شده استفاده می‌شوند. به عبارت دیگر در هر تکرار مؤلفه جدیدی به مدل آمیخته اضافه می‌شود به طوری که تابع لگاریتم درستنمایی (معیار اطلاع) مدل آمیخته جدید نسبت به تابع



شکل ۳: نمودار تابع چگالی مشاهدات و مدل آمیخته گاما با دو مؤلفه $F_2(x_t)$.

آمیخته فعلی اضافه می‌شود اگر

$$Q((1 - c_k^*)F_{k-1} + c_k^*f_k) - Q(F_{k-1}) > \frac{1}{\gamma} \chi_{1-\frac{\alpha}{\gamma}}^2(1).$$

به عبارت دیگر اضافه کردن مؤلفه جدید به مدل آمیخته فعلی باعث بهبود مدل برازش می‌شود. عملکرد دو روش یادگیری با استفاده از داده‌های شبیه‌سازی و مجموعه داده‌های واردات کالای ایالات متحده توسط گمرک ارزیابی شده است. با استفاده از مطالعات شبیه‌سازی نشان داده شده که ممکن است روش یادگیری آمیخته تقویت شده مدل درست را انتخاب نکند. اگرچه افزایش ناچیز تابع لگاریتم درست‌نمایی مدل آمیخته جدید باعث اضافه شدن مؤلفه جدید به مدل آمیخته فعلی می‌شود اما باعث بهبود مدل نمی‌شود.

آمیخته فعلی اضافه می‌شود که اختلاف تابع درست‌نمایی (معیار اطلاع) مدل آمیخته فعلی از مدل آمیخته جدید بزرگ‌تر (کوچک‌تر) از صفر باشد. اگر مقدار این اختلاف ناچیز و به صفر نزدیک باشد اضافه کردن مؤلفه جدید، مدل آمیخته فعلی را بهبود نمی‌بخشد. از این جهت این روش نیاز به تصحیح دارد. به عبارت دیگر روش یادگیری آمیخته تقویت شده با تابع هدف لگاریتم درست‌نمایی (معیار اطلاع) مدل آمیخته توانایی تشخیص مدل‌های معادل را ندارد. در این مقاله، روش یادگیری آمیخته تقویت شده بر اساس آزمون انتخاب مدل و ونگ تصحیح شده است. در روش آمیخته تقویتی تصحیح شده مؤلفه جدید با وزن بهینه c_k^* به مدل

مراجع

- [1] Du, J., Hu, Y., and Jiang, H. (2011). Boosted mixture learning of Gaussian mixture hidden markov models based on maximum likelihood for speech recognition. *IEEE TRANSACTIONS ON AUDIO, SPEECH, AND LANGUAGE PROCESSING*, **19**(7), 2091-2100.
- [2] Freund, Y., and Schapire, R. (1997). A decision theoretic generalization of on line learning and an application to boosting. *Journal of Computer and System Sciences*, **55** (1), 119-139.
- [3] Kim, M., and Pavlovic, V. (2007). A recursive method for discriminative mixture learning. *In Proc. ICML*, **24**, 409-416. <https://doi.org/10.1145/1273496.1273548>.
- [4] Kullback, S., and Leibler, R. A. (1951). Information and sufficiency modelling. *Annals of Mathematical Statistics*, **22**, 79-86.

- [5] Martins, I., Carvalho, P., Corte-Real, L., and Alba-Castro, J. L. (2018). BMOG: boosted Gaussian mixture model with controlled complexity for background subtraction. *Pattern Analysis & Applications*, **2(3)**, 641–654.
- [6] Mason, L., Baxter, J., Bartlett, P., and Frean, M. (1999). Boosting algorithms as gradient descent in function space. *NIPS*, **11**, 512-518.
- [7] McLachlan, G. (2001). *Finite Mixture Models*. New York: Wiley.
- [8] Ni, L., Wang, D., Wu, J., Wang, Y., Tao, Y., Zhang, J., and Liu, J. (2020). Streamflow forecasting using extreme gradient boosting model coupled with Gaussian mixture model. *Journal of Hydrology*, **586(1)**, 124901.
- [9] Pavlovic, V. (2004). Model based motion clustering using boosted mixture modeling. *In Proc. CVPR*, **1**, 811-818.
- [10] Schapire, R. (1990). The strength of weaklearnability. *Mach. Learn*, **5**, 197–227.
- [11] Schapire, R. (1999). Theoretical views of boosting and applications. *In: Watanabe, O., Yokomori, T. (eds) Algorithmic Learning Theory. ALT 1999. Lecture Notes in Computer Science, vol 1720*. Springer, Berlin, Heidelberg, 13-25.
- [12] Schneider, L., Chalmers, R. P., Debelak, R., and Merkle, E. C. (2020). Model selection of nested and non-nested item response models using Vuong tests. *Multivariate Behavioral Research*, **55(5)**, 664-684.
- [13] Tang, H., Chu, S., Hasegawa-Johnson, M., and Huang, T. (2009). Emotion recognition from speech via boosted Gaussian mixture models. *2009 IEEE International Conference on Multimedia and Expo*, New York, NY, US, 294-297.
- [14] Vlassis, N., and Likas, A. (2002). A greedy EM algorithm for Gaussian mixture learning. *Neural Process. Lett*, **15**, 77-87.
- [15] Vuong, Q. H. (1989). Likelihood ratio tests for model selection and non-nested hypotheses. *Econometrica*, **57**, 307-333.
- [16] Young, D. S., Chen, X., and Hewage, D. C. (2019). Finite mixture-of-Gamma distributions: Estimation, inference and model-based clustering. *Advances in Data Analysis and Classification*, **13**, 1053-1082.

Correcting Boosted Mixture Learning method using Vuong's test and its application in the Gamma Mixture Model

Sedigheh Zamani Mehreyan ¹

Abstract:

The boosted mixture learning method, BML, is an incremental method to learn mixture models for the classification problem. In each step of the boosted mixture learning method, a new component is added to the mixture model according to an objective function to ensure that the objective function is maximized. Sometimes the likelihood function or equivalently information criteria are defined as the objective function of BML. The mixture model is updated whenever a new component is added to the mixture model based on the maximum likelihood function and information criteria. Since the information criteria does not have the ability to identify equivalent models, therefore, it is possible that the new mixture model and the current mixture model are equivalent. In this paper, the boosted mixture learning method has been corrected using Vuong's model selection test, which has the ability to identify equivalent models. The performance of two learning methods is evaluated over simulation data and over the U.S. imports of goods by customs basis.

Keywords: Boosted mixture learning method, Maximum likelihood function, Model selection, Vuong's test.

¹ Department of Statistics, Imam Khomeini International University, Qazvin, Iran, zamani@sci.ikiu.ac.ir