

بسمه تعالی

دانشگاه

گزارش کارآموزی
آمار و سنجش آموزشی

عنوان

رگرسیون لجستیک

نگارش

استاد راهنما

سال ۹۴

فهرست

صفحه	عنوان
۳	مقدمه
۴	اهداف بررسی رگرسیون لجستیک:
۴	رگرسیون لجستیک چیست؟
۵	اهداف رگرسیون لجستیک:
۵	معرفی متغیر وابسته: binary
۵	منحنی لجستیک:
۶	پیش فرض های رگرسیون لجستیک
۶	تخمین مدل رگرسیون لجستیک:
۷	نحوه تعریف متغیرهای طبقه بندی شده (اسمی و ترتیبی) در رگرسیون لجستیک:
۸	انواع رگرسیون لجستیک
۱۰	رگرسیون لجستیک اسمی دووجهی
۱۰	پیش فرض ها:
۱۲	بیان مسأله:
۱۲	معرفی متغیرها:
۱۲	هدف این تحقیق:
۲۰	تحلیل داده ها
۲۵	نتیجه گیری:

مقدمه

همانطور که می دانیم برای انجام تحلیل رگرسیون خطی، متغیر وابسته باید کمی و در سطح سنجش فاصله ای / نسبی باشد. اما گاهی اوقات اتفاق می افتد که متغیر وابسته تحقیق در مقیاس فاصله ای / نسبی نبوده و مقیاس آن به صورت اسمی است. حال سوال اینجاست که برای این کار باید چه کرد. درحالی که پیش فرض اساسی تحلیل رگرسیون، مقیاس فاصله ای / نسبی متغیر وابسته است. در چنین حالتی نرم افزار SPSS این امکان را برای ما فراهم کرده است تا بتوانیم عوامل پیش بینی کننده تغییرات یک متغیر اسمی را نیز شناسایی کنیم. این روش که رگرسیون لجستیک نام دارد.

در اواخر دهه ۱۹۶۰ و اوایل ۱۹۷۰ به عنوان بدیلی برای روش رگرسیون خطی و همچنین تحلیل تابع تشخیصی مطرح شد.

زمانی که متغیر وابسته در سطح اسمی است و متغیرهای مستقل هم ترتیبی و هم فاصله ای هستند روشهای رگرسیون خطی معمولی و تحلیل تشخیصی مقدار برآورد ها را کمتر از مقدار واقعی نشان می دهند. رگرسیون لجستیک شبیه رگرسیون خطی است با این تفاوت که نحوه محاسبه ضرایب در این روش یکسان نمی باشد. به این معنی که رگرسیون لجستیک به جای حداقل کردن مجذور خطاها احتمالی را که یک واقعه روی میدهد حداکثر می کند.

همچنین در تحلیل رگرسیون خطی برای آزمون برازش مدل و معنی دار بودن اثر هر متغیر در مدل به ترتیب از آماره های F و t استفاده می شود در حالی که در لجستیک از آماره های کای اسکور و والد استفاده می شود.

رگرسیون لجستیک نسبت به تحلیل تشخیصی نیز ارجحیت دارد و مهمترین دلیل آن است که در تحلیل تشخیصی گاهی اوقات احتمال وقوع یک پدیده خارج از طیف صفر تا یک قرار می گیرد و متغیرهای پیش بین نیز باید دارای توزیع نرمال چند متغیره باشند در حالی که در رگرسیون لجستیک احتمال وقوع یک پدیده در داخل محدوده صفر تا یک قرار دارد و رعایت پیش فرض نرمال بودن متغیرهای پیش بین لازم نیست.

اهداف بررسی رگرسیون لجستیک:

1. توضیح دادن شرایطی که از رگرسیون لجستیک به جای رگرسیون چند گانه استفاده می کنیم.
2. معرفی نوع متغیرهای مستقل و متغیر وابسته در رگرسیون لجستیک
3. تفسیر نتایج آنالیز رگرسیون لجستیک و ارزیابی دقت پیش بینی
4. درک نقاط ضعف و قوت رگرسیون لجستیک در مقابل رگرسیون چند گانه و آنالیز تفکیک

رگرسیون لجستیک چیست؟

رگرسیون لجستیک و آنالیز تفکیک تکنیک های آماری مناسبی هستند زمانی که متغیر وابسته یک متغیر قیاسی (categorical) و متغیر های مستقل متریک یا غیر متریک باشند.

رگرسیون لجستیک حالت خاصی از رگرسیون می باشد که برای پیش بینی و توضیح یک متغیر زوجی (binary) فرمول بندی شده است.

رگرسیون لجستیک از این لحاظ که فقط متغیرهای زوجی را پیش بینی می کند یک رگرسیون محدود شده می باشد.

بنابراین زمانی که گروه ها 3 یا بیشتر از 3 هستند آنالیز تفکیک (discriminant analysis) نتیجه بهتری را به ما بدست می دهد. رگرسیون لجستیک یک متغیر وابسته زوجی را در فرم عمومی به شکل زیر براساس مجموعه ای از متغیر های مستقل پیش بینی می کند:

$$Y_1 = X_1 + X_2 + X_3 + \dots + X_n$$

(binary nonmetric) (nonmetric and metric)

لجستیک رگرسیون به دو دلیل به آنالیز تفکیک ترجیح داده شود:

1. آنالیز تفکیک به شدت به برقراری پیش فرض های (نرمال بودن و تساوی واریانس-کوواریانس) حساس می باشد در حالی که در اکثر مواقع این پیش فرض ها برقرار نبوده و رگرسیون لجستیک robust می باشد.
2. حتی اگر پیش فرض ها نیز برقرار باشد محققان ترجیح می دهند که به دلیل شباهت خیلی زیاد رگرسیون لجستیک به رگرسیون چند گانه از رگرسیون لجستیک استفاده کنند.

اهداف رگرسیون لجستیک:

1. تشخیص متغیرهای مستقلی که بر عضویت در گروه برای متغیر وابسته اثر می گذارند.
2. ساخت سیستم طبقه بندی شده براساس مدل لجستیک برای تصمیم گرفتن در مورد عضویت گروه.

رگرسیون لجستیک چندین مشخصه منحصر به فرد دارد که بر روی طرح تحقیقی اثر دارد:

1. اولین خصیصه، طبیعت منحصر به فرد متغیر وابسته است که آن زوجی (binary) بودن متغیر می باشد.
2. دومین خصیصه مربوط می شود به اندازه نمونه، که چند عامل بر روی این خصیصه اثر گذار می باشند، از جمله روش تخمین استفاده شده در رگرسیون لجستیک (MLE)

معرفی متغیر وابسته: binary

رگرسیون لجستیک برای متغیر زوجی وابسته دو گروه از علایق را در نظر میگیرد، که به یکی از گروه ها مقدار 0 و به گروه دیگر مقدار 1 را اختصاص می دهد و این اصلاً مهم نیست که کدام مقدار 1 (یا 0) به کدام گروه تخصیص داده می شود. اما در زمان تفسیر ضرایب باید این موضوع مورد توجه قرار بگیرد.

منحنی لجستیک:

به دلیل اینکه متغیر وابسته زوجی (binary) تنها مقادیر 1 و 0 را می گیرد، مقادیر پیش بینی شده (احتمال ها) برای اینکه در محدوده های یکسانی قرار گیرند، باید محدوده بندی شوند. برای تعریف یک رابطه محدوده بندی شده بوسیله 1 و 0، رگرسیون لجستیک از منحنی لجستیک برای نشان دادن رابطه بین متغیر وابسته و متغیر مستقل استفاده می کند.

استفاده از منحنی لجستیک این واقعیت را نشان میدهد اگر حتی در یک رابطه رگرسیون اصلاحیه هایی برای تبدیل اثرهای غیر خطی به خطی هم انجا میشود نمی تواند این موضوع را که مقادیر پیش بینی شده بین 1 و 0 باقی بماند را تضمین کند.

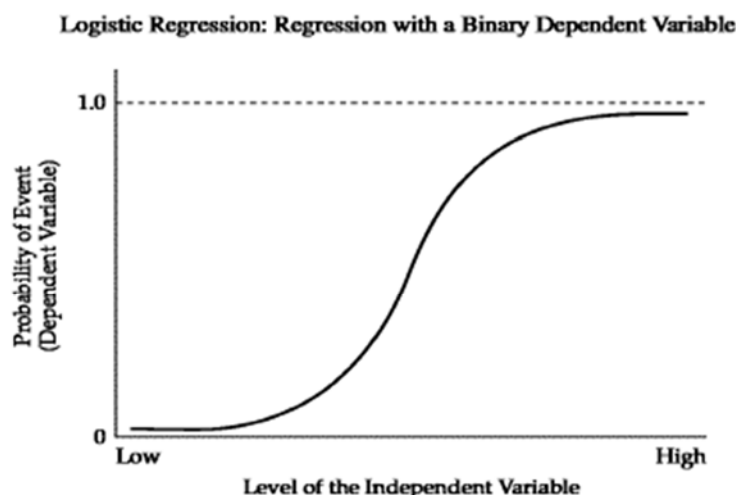


FIGURE 1 Form of the Logistic Relationship Between Dependent and Independent Variables

پیش فرض های رگرسیون لجستیک

مهم ترین پیش فرض رگرسیون لجستیک زوجی بودن متغیر وابسته می باشد. در رگرسیون لجستیک نیازی نیست که متغیر های مستقل از توزیع خاصی پیروی کنند و متریک یا غیر متریک باشند همچنین در رگرسیون لجستیک نیازی به این نیست که رابطه بین ضرایب متغیر وابسته و ضرایب متغیر های مستقل یک رابطه خطی باشد در صورتی که در رگرسیون چندگانه به این صورت نبود.

تخمین مدل رگرسیون لجستیک:

رگرسیون لجستیک همانند رگرسیون چند گانه دارای یک متغیر وابسته است که تشکیل شده از ضرایب پیش بینی شده ای برای هر یک از متغیر های مستقل با این تفاوت که به شیوه ای متفاوت نسبت به

رگرسیون چند گانه تخمین زده می شود. رگرسیون لجستیک اسمش را از تبدیلات لجیتی (logit transformation) گرفته است.

نحوه تعریف متغیرهای طبقه بندی شده (اسمی و ترتیبی) در رگرسیون لجستیک:

یکی از مهمترین مشکلات در اجرای تحلیل رگرسیون لجستیک وجود متغیرهای ترتیبی است. در هنگام اجرای رگرسیون لجستیک فرض بر این است که تمامی متغیرهای مستقل در سطح سنجش فاصله ای/نسبی هستند در حالی که در عمل چنین نیست و برخی از آنها اسمی و ترتیبی نیز هستند اما از آنجا که رگرسیون لجستیک با نسبت احتمال وقوع یک پدیده به احتمال عدم وقوع آن پدیده سر و کار دارد بنابر این متغیرهای مستقل حتما باید به متغیرهای شبه فاصله ای با در کد 1 و 0 تبدیل شوند تا بتوانیم نسبت طبقات آن در متغیر وابسته را بررسی کنیم به همین خاطر در نرم افزار SPSS در هنگام اجرای دستور رگرسیون لجستیک از طریق کادر Categorical... در کادر اصلی دستور این امکان وجود دارد که متغیرهای طبقه بندی شده (اسمی و ترتیبی) را به صورت مجازی به متغیرهای فاصله ای تبدیل کنیم.

برای مجازی کردن متغیرهای اسمی و ترتیبی باید هر یک از طبقات آن متغیر به عنوان یک متغیر جداگانه با دو طبقه تعریف شده و به طبقه اول کد 1 و به طبقه دوم کد 0 تعلق گیرد.

به عنوان مثال اگر متغیر مورد نظر ما سطح تحصیلات است که در طبقات پایین، متوسط و بالا تعریف شده است باید هر گزینه را به عنوان یک متغیر دو وجهی حساب کرده و به کسانی که آن میزان تحصیلات را دارند کد 0 و به کسانی که آن میزان تحصیلات را ندارند کد 1 تعلق گیرد.

یعنی بدین صورت که (متغیر اول) تحصیلات پایین مساوی یک و تحصیلات غیر پایین مساوی صفر.

(متغیر دوم) تحصیلات متوسط مساوی یک و تحصیلات غیر متوسط مساوی صفر.

نکته: 1 همانطور که در طبقه بندی بالا ملاحظه می شود متغیر تحصیلات در هنگام تبدیل متغیر مجازی فقط در دو طبقه تعریف شده و طبقه سوم حذف شده است.

دلیل این امر آن است که در رگرسیون لجستیک همانند رگرسیون خطی متغیر مجازی برای طبقه آخر تعریف نمی شود و تعداد آن همواره باید یکی کمتر از تعداد طبقات متغیر اصلی باشد این اصل برای اجتناب از مسئله تکنیکی چند هم خطی بودن در رگرسیون لجستیک است.

طبقه ای که متغیر مجازی تبدیل نمیشود طبقه مرجع نام دارد که مبنای مقایسه و تقابل با سایر طبقات قرار می گیرد.

نکته : 2 موقعی که طبقات متغیر مستقل با طبقات متغیر وابسته به منظور مقایسه در تقابل قرار می گیرند در هنگام اجرای کادر Categorical... در دستور رگرسیون لجستیک امکان انتخاب چندین نوع تقابل وجود دارد:

1- شاخص: در این روش تقابل ها به صورت عضویت یا عدم عضویت در یک طبقه نشان داده میشود . طبقه مرجع نیز به صورت یک ردیف در ماتریس تقابل با مقادیر (۱) نشان داده می شود. این روش رایج ترین روش انتخاب تقابل ها است که اغلب نیز از این روش استفاده میکنند.

2- ساده: در این روش هر طبقه از متغیر پیش بین با طبقه مرجع با طبقه مرجع متغیر وابسته مقایسه می شوند.

3- تفاوت: هر طبقه از متغیر پیش بین با میانگین اثر طبقات قبلی مقایسه میشود این روش به معکوس تقابل های هلمرت نیز معروف است.

4- هلمرت: هر طبقه از متغیر پیش بین با میانگین اثر طبقات بعدی مقایسه می شوند.

۵- چند جمله ای: در این روش که به تقابل های چند جمله ای متعامد معروف است فرض بر این است که فاصله بین طبقات برابر باشد. این تقابل ها فقط برای متغیر های عددی امکان پذیر هستند.

۶- انحراف: هر طبقه از متغیر های پیش بین با اثر کل مقایسه می شود.

انواع رگرسیون لجستیک

همان طور که در ابتدای مبحث تحلیل رگرسیون لجستیک گفته شد در رگرسیون لجستیک متغیر وابسته می تواند به دو شکل دو وجهی و چند وجهی باشد. به همین خاطر در نرم افزار SPSS شاهد وجود دو نوع تحلیل رگرسیون لجستیک هستیم که بسته به تعداد مقولات و طبقات متغیر وابسته می توانیم از یکی از این دو شکل استفاده کنیم:

1- رگرسیون لجستیک اسمی دو وجهی: موقعی است که متغیر وابسته در سطح اسمی دو وجهی است. یعنی زمانی که با یک متغیر وابسته اسمی دو وجهی سر و کار داریم.

2-رگرسیون لجستیک اسمی چند وجهی یا چند جمله ای: موقعی مورد استفاده قرار می گیرد که متغیر وابسته اسمی چند وجهی باشد.

روش های انتخاب متغیر ها در رگرسیون لجستیک در رگرسیون لجستیک روش های متعددی برای انتخاب و ورود متغیر ها به مدل وجود دارد که به ما کمک میکنند تا مشخص کنیم که چگونه متغیر های مستقل وارد تحلیل شوند و نیز بتوانیم مدل های رگرسیونی متفاوتی را بر روی یک مجموعه متغیر یکسان ایجاد کنیم.

1-روش همزمان:در این روش تمامی متغیرها در یک مرحله وارد مدل می شوند.

2-روش پیشرو مشروط: نوعی روش گام به گام است که در آن ورود متغیرها به تحلیل بر اساس معنی داری مقدار آماره نسبت درست نمایی و خروج متغیرها از تحلیل بر اساس احتمال این آماره و با توجه به بر آورد پارامتر مشروط انجام میگردد.

3-روش پیش رو نسبت درست نمایی:نوعی روش گام به گام است که در آن ورود متغیر ها به تحلیل بر اساس معنا داری مقدار آماره نسبت درست نمایی و خروج متغیرها از تحلیل بر اساس احتمال این آماره و با توجه به برآوردهای حداکثر درست نمایی جزئی انجام میشود.

4-روش پیش رو والد:نوعی روش گام به گام است که در آن ورود متغیر ها به تحلیل بر اساس معنی داری مقدار آماره نسبت درست نمایی و خروج متغیر ها از تحلیل بر اساس احتمال آمار والد انجام میگردد.

5-روش حذف پس رو مشروط: نوعی روش گام به گام پسرو است که در آن خروج متغیر ها از تحلیلی بر اساس احتمال آماره نسبت درست نمایی و با توجه به برآوردهای پارامتر مشروط انجام می گیرد.

6-روش حذف پسرو نسبت درست نمایی: نوعی روش گام به گام پسرو است که در آن خروج متغیرها از تحلیل بر اساس احتمال آماره نسبت درست نمایی و با توجه به برآورد حداکثر درست نمایی جزیی انجام می گیرد.

7-روش حذف پسرو والد:نوعی روش گام به گام پسرو است که در آن خروج متغیرها از تحلیل بر اساس احتمال والد انجام می گیرد.

رگرسیون لجستیک اسمی دووجهی

تحلیل رگرسیون لجستیک اسمی دو وجهی زمانی مورد استفاده قرار می گیرد که متغیر وابسته در سطح اسمی دو وجهی است و بنا داریم وجود یا عدم یک صفت را بر اساس مجموعه ای از متغیرهای مستقل پیش بینی می کنیم.

بنابر این در رگرسیون لجستیک اسمی دووجهی ما نمی توانیم مانند رگرسیون خطی چندمتغیره مقدار عددی دقیق یک متغیر وابسته را بر اساس اطلاعاتی که راجع به متغیرهای مستقل داریم تعیین کنیم بلکه در این روش ما با نسبت احتمال سر و کار داریم که آن را با کد 0 نشان می دهند تا کد 1 .

مثال های متعددی را می توان بر شمرد که در آن با یک متغیر وابسته اسمی دو وجهی سر و کار داریم این که چرا برخی کودکان در هنگام تولد می میرند و برخی دیگر زنده می مانند؟

این که چرا برخی از شهروندان در انتخابات شرکت می کنند و برخی شرکت نمی کنند؟ و غیره .

این ها همه سوالاتی هستند که نرم افزار SPSS این امکان را برای ما فراهم می کند تا بتوانیم عوامل پیش بینی کننده تغییرات یک متغیر اسمی دو وجهی را نیز شناسایی کنیم.

در واقع در این نوع رگرسیون لجستیک می توانیم احتمال وقوع یک واقعه را به طور مستقیم بر آورد کنیم.

پیش فرض ها:

- 1-متغیر وابسته حتما باید در سطح سنجش اسمی دو وجهی باشد.
- 2-متغیر های مستقل می توانند هم در سطح کمی و هم در سطح کیفی طبقه بندی شده باشند . اما چنانچه یک یا چند متغیر مستقل در سطح اسمی/ترتیبی بودند حتما باید ابتدا این متغیرها را به متغیر های مجازی تبدیل کنیم (یعنی 1 و 0) البته در رگرسیون لجستیک ، کادری به نام Categorical... وجود دارد که با انتخاب و اجرای آن متغیر های ترتیبی به طور خودکار به متغیر های مجازی تبدیل می شوند بنابراین نیازی به کد گذاری مجدد آنها توسط محقق نیست.
- 3-لزوم تبعیت داده های متغیر های مستقل از توزیع نرمال ضروری نیست. اما چنانچه این متغیرها دارای توزیع نرمال چندمتغیره باشند ، در آن صورت برآزش مدل بهتر خواهد بود.

4-چند هم خطی نبودن متغیر های مستقل از دیگر مفروضات رگرسیون لجستیک است. چرا که در صورت چند هم خطی بودن این متغیرها برآوردها دارای اریب بوده و خطاهای استاندارد نیز نوسان زیادی خواهند داشت.

ترسیم نمودار پراکنش به ما کمک می کند تا از چند هم خطی بودن یا نبودن متغیر های مستقل اطمینان حاصل کنیم.

نکته: چنانچه پیش فرض های نرمال بودن چند متغیره و برابری ماتریس های واریانس و کوواریانس تامین شدند در آن صورت پیشنهاد میشود که از روش تحلیل لجستیک استفاده کنیم.

بیان مسأله :

آیا عواملی مانند معدل، منطقه محل زندگی (از نظر محرومیت)، میزان تحصیلات، سن و... در آینده شغلی یک فرد مؤثر است یا خیر؟

این تحقیق به بررسی عوامل (سن-جنسیت-معدل دیپلم-معدل آخرین مدرک تحصیلی-میزان تحصیلات-محل زندگی) بر داشتن یا نداشتن شغل با استفاده از رگرسیون لجستیک میپردازد.

همانطور که در بالا گفته شد برای بررسی با رگرسیون لجستیک نیاز به یک متغیر وابسته به صورت ۱ و ۰ و چند متغیر مستقل داریم.

نمونه گیری روی ۴۰ نفر از افراد استان تهران (بامناطق مختلف)-۲۰ مرد و ۲۰ زن با میزان تحصیلات مختلف و سن بین ۲۰ تا ۴۰ سال بارشته های متفاوت انجام گرفته است.

معرفی متغیرها:

متغیر وابسته (y) که در اینجا :داشتن یا نداشتن شغل است.

متغیرهای مستقل (x) که در اینجا:

سن-جنسیت-معدل دیپلم-معدل آخرین مدرک تحصیلی-میزان تحصیلات-محل زندگی است.

سن: ۲۰ تا ۴۰

جنسیت: مرد یا زن

میزان تحصیلات: دیپلم-لیسانس-فوق لیسانس-دکتری

محل زندگی: محروم یا غیر محروم

رشته: ریاضی-انسانی-تجربی-هنر

هدف این تحقیق:

هدف ما از انجام این تحقیق این است که ببینیم آیا سن یک فرد (زن یا مرد)-معدل او-محل زندگی و میزان تحصیلاتش باعث میشود که فرد شغل داشته باشد؟ یاخیر این عوامل به داشتن شغل هیچ ارتباطی ندارند؟

اگر متغیرهای مستقل ما در ایجاد شغل نقش داشتند چه کارهایی باید انجام دهیم و اگر این متغیرها در ایجاد شغل نقشی ندارند عوامل مؤثر چه عواملی هستند؟

فرض ما براین است که افرادی که مدرک تحصیلیشان بالاتر است و معدل بیشتری دارند و در مناطقی که از نظر امکانات محروم نیستند دارای شغل هستند.

وضعیت شغل	رشته	معدل مدرک	معدل دیپلم	منطقه	مقطع	سن	جنسیت	افراد
۰	۱	۱۶	۱۶,۵	۲	۲	۲۲	۰	۱
۰	۲	۱۱,۵	۱۲,۶۱	۳	۱	۳۰	۱	۲
۱	۱	۱۸,۳۰	۱۸,۴۱	۱	۲	۲۸	۱	۳
۰	۳	۱۹	۱۹	۲	۲	۲۳	۱	۴
۰	۴	۱۸,۱۴	۱۸,۷۵	۱	۳	۲۵	۰	۵
۱	۱	۱۵	۱۷,۲۹	۳	۴	۳۲	۰	۶
۱	۲	۱۵	۱۶,۹۰	۲	۴	۳۲	۰	۷
۱	۴	۱۲	۱۲,۸۰	۳	۱	۳۰	۱	۸
۰	۱	۱۴,۳۳	۱۵,۴۰	۱	۱	۲۰	۰	۹
۰	۳	۱۳,۱۲	۱۴	۳	۲	۲۴	۰	۱۰
۱	۴	۱۹,۲۰	۱۹,۳۴	۱	۳	۲۸	۱	۱۱
۱	۲	۱۲,۸۵	۱۳,۴۶	۲	۱	۲۵	۱	۱۲
۰	۴	۱۷,۲۲	۱۷,۴۰	۳	۳	۳۳	۰	۱۳
۱	۱	۱۰,۹۳	۱۱,۲۰	۱	۱	۳۶	۰	۱۴
۰	۳	۱۴,۳۸	۱۴,۸۶	۱	۲	۲۷	۱	۱۵
۰	۴	۱۹,۷۳	۱۹,۹۰	۳	۳	۲۹	۰	۱۶
۰	۴	۱۹	۱۸,۸۸	۲	۱	۲۱	۱	۱۷
۱	۱	۱۴,۴۱	۱۵,۱۸	۱	۲	۳۰	۱	۱۸
۱	۲	۱۶,۱۲	۱۶,۲۰	۲	۴	۳۴	۰	۱۹
۱	۳	۱۹,۳۰	۱۹,۴۸	۳	۴	۳۶	۰	۲۰

وضعیت شغل	رشته	معدل مدرک	معدل دیپلم	منطقه	مقطع	سن	جنسیت	افراد
۱	۴	۱۷	۱۷	۲	۲	۲۵	۱	۲۱
۰	۳	۱۴	۱۴,۱۲	۱	۱	۲۹	۱	۲۲
۱	۲	۱۴	۱۵,۶۰	۳	۴	۳۲	۱	۲۳
۰	۱	۱۲,۲۰	۱۴,۳۰	۲	۳	۳۱	۰	۲۴
۰	۳	۱۱,۴۵	۱۱,۸۰	۲	۱	۳۱	۱	۲۵
۰	۲	۱۲	۱۲,۲۱	۳	۱	۲۰	۰	۲۶
۰	۴	۱۳,۸۰	۱۴,۸۷	۱	۲	۲۲	۰	۲۷
۰	۱	۱۴	۱۴,۳۰	۳	۲	۲۶	۱	۲۸
۰	۲	۱۴,۸۱	۱۵,۹۰	۱	۳	۳۲	۰	۲۹
۱	۴	۱۹,۹۰	۱۹,۹۶	۲	۴	۳۵	۱	۳۰
۱	۲	۱۶	۱۶	۳	۳	۲۸	۰	۳۱
۱	۳	۱۸	۱۹	۱	۳	۳۱	۱	۳۲
۱	۲	۱۸,۷۲	۱۷,۹۰	۱	۳	۲۹	۰	۳۳
۰	۲	۱۴,۲۰	۱۴,۸۴	۳	۱	۲۳	۱	۳۴
۰	۳	۱۹,۱۲	۱۹,۴۰	۳	۴	۳۵	۰	۳۵
۰	۲	۱۸	۱۸,۵۰	۱	۲	۲۴	۰	۳۶
۱	۳	۱۴	۱۵	۳	۴	۳۱	۱	۳۷
۱	۳	۱۶,۳۲	۱۷,۸۸	۲	۳	۲۶	۱	۳۸
۱	۱	۱۶,۵۶	۱۶,۸۷	۲	۴	۳۳	۰	۳۹
۱	۲	۱۸,۹۵	۱۸,۹۴	۲	۴	۳۴	۱	۴۰

Logistic Regression

Case Processing Summary

Unweighted Cases ^a		N	Percent
Selected Cases	Included in Analysis	40	100.0
	Missing Cases	0	.0
	Total	40	100.0
Unselected Cases		0	.0
Total		40	100.0

a. If weight is in effect, see classification table for the total number of cases.

**Dependent Variable
Encoding**

Original Value	Internal Value
bikar	0
shaghel	1